

Machina Cogitans: The Promethean Fire Yet Smoulders

T. J. Hayes

Toronto, Ontario*

Abstract. Many ancient cultures share fire myths that frame fire as the foundational tool that allowed humans to prosper yet also carries the potential for destruction, a metaphor for humanity leaving the Savanna for civilization. Mumford suggests language was the first tool that allowed all others to be conceived, yet leaves no trace in the fossil record. Language is the core of a new technology, Artificial Intelligence, which is placing humanity on a transformative new course. Using McGilchrist’s hemispheric hypothesis and Vervaeke’s Four Ways of Knowing, we propose the limits of AI to be bounded to left hemispheric cognition, Propositional and Procedural Knowing, while embodied, Participatory knowing remains ontologically beyond reach. We acknowledge this via a new taxonomic domain, *Silicata*, with AI named *Machina Cogitans*: a thinking machine, not a wise one. Doing so invites research across disciplines, provides a new framing for humans to interact with *Machina Cogitans*, and suggests the need for a pedagogy emphasizing right hemispheric attention, ensuring we co-evolve with AI as discerning, moral agents rather than passive consumers of machine cognition.

Keywords: Fire Myths, Language, Artificial Intelligence, Participatory Knowing, Hemispheric Lateralization, *Machina Cogitans*

“Prometheus stole fire from the gods. We are each the heirs of that divine spark. Used wisely, the spark fuels one’s journey and lights the way. Treated carelessly, the spark consumes its owner and everything in its path.”

— Thomas Lloyd Qualls

Fire Myths

ALMOST EVERY CULTURE HAS a myth of stealing fire from the gods and its relevant lesson: powerful tools can greatly benefit society, but used recklessly can lead to chaos, disorder, and destruction. In each story, the society realizes, that progress always demands a price, and that power without wisdom can lead to calamity or complete ruin. Further, in each of these stories, a price is always paid and scar tissue remains, be it fracturing the cosmic order of the world, or the punishment of the gods to restrain humans, or some personal tragedy. In the story of Prometheus, Pandora is sent to mankind via Epimetheus, and Zeus’ retribution on humanity is delivered once Pandora opens the jar, releasing disease, toil, despair, suffering, and other evils into the world. Where Prometheus is the *fore-thought* of hope and possibilities fire represents for humanity, Epimetheus is the *after-thought*, the consequence that follows.

Of course, these myths are not literal stories about how mankind obtained fire for the first time. Rather, they are deeper stories about who we are, and contain lessons for those who follow after. As [Schomann](#) explains, “myth functions as a metanarrative — an organizing story that confers coherence, intelligibility, and moral orientation, regardless of its empirical status.” In the Prometheus story fire is knowledge, technology, and civilization; the light of consciousness itself. Yet, knowledge has the power both to liberate or cause suffering. Regardless of which form it takes, it always changes the wielder. There is no return to Eden.

Man as Tool Maker

THE FACT THAT MASTERY-OF-FIRE stories are often foundational in the earliest myths of many cultures is telling, in that it lends currency to the conceit that humanity’s place in nature, at the top, is because of our tool-making ability. We read about the pivot location of our thumbs allowing us to manipulate objects with more dexterity than other apes; we refer to parts of the world as ‘developed’ or ‘emerging,’ with respect to their technology; we define epochs in history at markers of technological prowess, such as the Stone Age, the ‘Bronze Age,’ the ‘Copper Age,’ the ‘Iron Age,’ and so forth.¹ That is, we classify epochs by the sophistication and materials of the tools left behind in the archaeological record. Ancient cultures asserted their mastery by producing megaliths such as the Pyramids of Giza, Stonehenge, and *Teōtihuacān*; conversely, modern society seeks to master the atomic scale. We interpret these as important waypoints on our march of progress toward mastery of nature. Tools separate man from nature, man from man.



Figure 1. Stanley Kubrick’s iconic match cut in a side-by-side view of the transition from a bone to orbiting station, representing the continuity of Man’s use of tools, and tools as weapons, spanning over 4,000,000 years. [[Sherlock, 2021](#)]

* Affiliated with The Apocrypha [Substack](#) and [Medium](#) pages.

In a more modern telling, Kubrick's *2001: A Space Odyssey* explores humanity's development via the intervention of an external (potentially alien) monolith. Our progression of tool mastery is visually represented by Kubrick in stunning form during an early scene, as we transition from bone tool to satellite station,² and is shown in Figure 1.

In viewing the world this way, we come to see tools as things, tangible and manipulable. In some respects, this is our inheritance from the Enlightenment, where Copernicus halted the wanderings of the Sun in the sky; Bacon's *Novum Organum* challenged Aristotle's thousand-year reign of philosophical reasoning with a new epistemology, the Scientific Method; and Newton stood on the shoulders of giants to see further than any before him, gifting the world the laws of physics. Using these tools, we have visited the Moon, enhanced medicine, and harnessed the very power of the stars, a testament to their power. In telling ourselves this modern story, over and over again, that we are the apes who mastered tools on our way to mastering nature, we come to see ourselves in the glow of this brighter light cast by Prometheus' fire as *homo faber*.

But Mumford [1967] challenges this framing. He points out the oversight, so to speak, that this myopic view is blind to core human abilities, abstraction and articulation, which are not laid down in the deep archaeological record: language is the *ur-tool*. We are more than clever tool makers, and limiting ourselves to such a view is incomplete. With language, we gain the ability to engage reality for continuity and coordination amongst others; we create meaning through myth and story; we pass on wisdom and lessons learnt to future generations; we manipulate concepts and anticipate the future by abstracting ideas into symbols. Before we could have tools that demonstrate manipulation, intention, and exhibit consistency over time³, language had to be there first to share this knowledge. As language becomes standardized over a region, stories, rituals, and myths become shared; civilization itself becomes possible.

Hemispheric Lateralization

 *The Master and His Emissary* and, more thoroughly, in *The Matter With Things*, McGilchrist [2009, 2021] sheds further light on this notion of seeing the world as "things" to be decomposed, itemized, and thus mastered by us through our tools. In and of itself, this is not a problem; it becomes pathological when that stance toward the world dominates. While his work spans a broad range of implications for society, philosophy, art, and meaning, well beyond the scope of this essay, it is worth a brief digression into his hemispheric hypothesis to help contextualize the arguments presented here.⁴

Unless you have read McGilchrist, be prepared to challenge everything you know, or think you know, about the differences between the hemispheres of the brain, their function, how they relate to each other, and, more importantly, how they each interpret the world. We are born with two hemispheres for a reason: attention. It is how the different hemispheres see and attend to the world that matters most. Attention, in this telling, should not be confused with 'focus,' but rather understood as each hemisphere having a fundamentally differing mode of *attending* reality, from which its interpretation follows.

In what is now a standard summary,⁵ consider a bird that needs to search the ground carefully to find a seed to eat amongst some pebbles. To do this, the bird needs to focus its attention, differentiate between seeds and pebbles, and then, once found, grab the seed with its beak to eat. However, if all the bird did was focus on seeds versus pebbles, it would soon become food for a predator. As such, the bird must have both a way of being able to sustain

narrow, goal-oriented attention and, *simultaneously*, broad, vigilant awareness. According to McGilchrist, these modes of attention are lateralized to the left and right hemispheres respectively.⁶

In his analysis, McGilchrist walks us through his research exploring how the two hemispheres differ from each other. The left hemisphere is superior in many ways at categorizing, labelling, modelling, and ultimately apprehending the world. Whereas the right hemisphere sees the world as it is, understands context, the *gestalt*, and is open to possibilities. In other words, it comprehends the world. This way of attending to the world allows the right hemisphere to *integrate* the information from the left hemisphere into a more coherent whole, from which wisdom can arise. The left hemisphere, for all its intellect, does not know what it does not know; it exhibits certainty and will confabulate to ensure the consistency of its model of the world. This is not to diminish the essential role the left hemisphere plays, but rather to highlight the limits necessarily imposed by its superiority in other vital aspects of understanding.

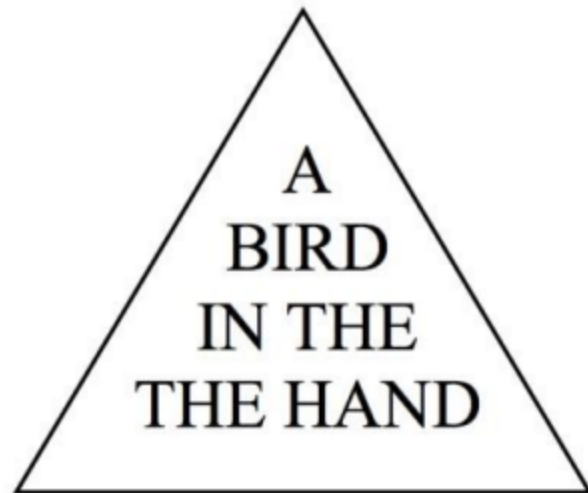


Figure 2. A test where subjects are to read the above figure in which the word 'the' is repeated.

In one of my favourite demonstrations in which context matters, McGilchrist provides this example, as shown in Figure 2 (which he, in turn, borrows from Kay [2011]). "John Kay writes, of those who miss the repetition: 'This is a mistake — or so the experimenter thinks. But who is really making the mistake: the pedant who offers the direct answer and reads 'a bird in the the hand'? Or the person with a life, who approaches the problem obliquely and valiantly finds sense in nonsense?'" [McGilchrist, 2021, Chapter 18].

These differing ways of attending to the world have implications not only for how we interpret reality, but for how we discover truth. What is the 'correct' answer in Figure 2? Well, it depends. Indeed, oftentimes the deepest truths are paradoxes that are true in both hemispheres, yet elude pure articulation. For example, McGilchrist often cites the *Dao De Jing*: "The dao that can be named is not the dao." When functioning in harmony, each mode of attention complements the other. But, as implied in the title of *The Master and His Emissary*, the right hemisphere occupies the role of Master, the left of Emissary. In the right, we have metaphor, context, and integrated *wisdom* through which

meaning coheres; in the left, incomplete, abstracted, calculated *thinking* to manipulate our environment.

Bringing it back to Mumford, I suggest the left hemisphere would be more likely to identify itself with *homo faber*, and then proceed to build ever more useful and powerful tools, sometimes for the sake of building them alone.

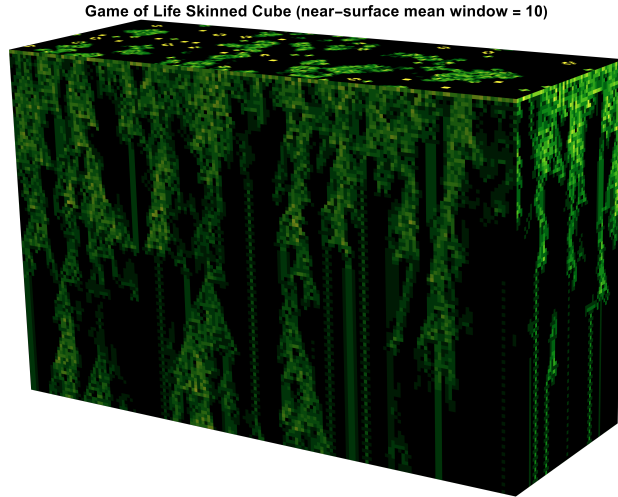


Figure 3. Showing near-surface averaged *Game of Life* Cube with near-surface averaging of $n = 10$. Block dimensions are: $x = 160$, $y = 60$, and $t = 100$, where time (t) increases from the top of the block to the bottom. Made with *Mathematica*.

The Last Invention

IN 1965, IRVING JOHN GOOD wrote, “*The first ultraintelligent machine is the last invention that man need make,*” **Good** [1965]. When he wrote that in 1965, the reality of creating ‘ultra’ or ‘superintelligence,’ as **Bostrom** famously popularized it, seemed necessary for the survival of humanity. Superintelligence, however, was still science fiction over half a century ago.⁷ Fast-forward to November 30, 2022: OpenAI released ChatGPT 3.5 to the public, or, as some would say, introduced us to Artificial Intelligence (AI). *A first contact*. What I seek to focus on is clarifying what *kind* of intelligence we are talking about.

In discussing machine intelligence, it helps to start with the simplest models. There is a field of study that explores simple rule-based systems to examine how complex patterns emerge. **Conway’s Game of Life** is one such laboratory, often used to show how even very simple rules can lead to complex, lifelike behaviour and emergent, organic landscapes, as shown in **Figure 3**. While no one would suggest that the cells in the Game of Life approach anything related to higher intelligence, they nonetheless exhibit behaviour that mimics simple forms of life. So the question becomes: what behaviour does AI currently mimic? The answer, of course, is found in the term employed as their more general name: Large Language Models (LLMs).

Reprise: Language as a Tool

HERE WE ARE AGAIN, back to Mumford. Recall that language was a necessary scaffold, a precondition, from which other tools come to be built. It is worth emphasizing how unique a tool language is. For starters, it is something

that one does not typically ‘see’ in the sense that its locus is found in the mind. Here, I am making the distinction between language *as such* and language as it is *re-presented* in written, spoken, or signed forms; that is, the external forms through which a relationship between the speaker and the spoken-to is established.

Language holds deep meaning for us, not only in communicating ideas, but also in establishing reality. For example, someone is not guilty until they are *pronounced* guilty by someone with the power of those words; we yield power to a priest or an officer of the court to *pronounce* a couple “husband and wife.” Language can also go horribly and tragically wrong when artificially promoted as political watchwords or when a population is propagandized, as described in **Klemperer** [2006]. Indeed, part of the disquiet one feels when reading Orwell’s *1984* is, in no small measure, the ‘Newspeak’ employed by Ingsoc to control the population by constraining the range of thought itself. But language can also be the vehicle of creation, as found in the sonnets and plays of Shakespeare, the poetry of Shelley, the *Epic of Gilgamesh*, and so on.

It should come as no surprise to anyone, then, that the most powerful tool humanity is creating is, at its foundation, employing language in all its forms to both learn about and interact with the world. We provide written instructions to LLMs; we chat with them, prompt them. And they, in turn, communicate back to us or carry out a range of instructions via computer language, and so on. By employing the *ur-tool* of language, Large Language Models may, in fact, become the *all-tool*: the Last Invention of Good.

Four Ways of Knowing

COGNITIVE SCIENTIST AND RESEARCHER John Vervaeke brings yet another dimension into the discussion of intelligence by providing a framework of knowing with increasing levels of depth. Known as the *Four “P”s*, this framework provides a useful matrix for assessing the capabilities of AI in terms of knowing. Moreover, his broader metacognitive theory emphasizes the need to integrate these forms of knowing and to evaluate their relevance within any given situation, which is how wisdom and meaning are found.

In the three-volume series *Mentoring the Machines*, co-written with Shawn Coyne [**Vervaeke and Coyne**, 2023–2025], the authors explore the ontological implications of bringing forth a new form of intelligence and our collective responsibility in ushering it forward. Indeed, one of the motivating insights for this essay arose from those works: how should we think about intelligence in the first place, given the new reality AI presents to society? More interestingly, is it possible to help Artificial General Intelligence (as distinct from narrow, human-objective-set AI) co-evolve alongside humanity under human mentorship?⁸

Knowing, within Vervaeke’s framework, is constituted by four increasing levels of attainment, from the most basic and accessible to the most deeply embodied way of knowing we have: Propositional, Procedural, Perspectival, Participatory. Propositional knowing is theoretical knowledge and widely accessible. This would be like reading a book on swimming; it becomes Procedural knowing when one actually jumps in the lake. Perspectival knowing is where discernment starts to enter the frame. Here, the ability to discern what is salient in the moment, from a particular perspective, is core to his concept. It represents the ability to begin interpreting the world through a subjective lens, not only your own, but that of others as well. Participatory knowing,

the final stage in this framework, is both experiential and recursive, in that participation itself becomes the way in which knowledge is embodied: this is the space where wisdom and identity are molded.

Another theme core to Vervaeke's framework is what he terms *relevance realization*, the ongoing, embodied process by which a knower discerns what matters in a given situation. The ability to orient oneself according to embodied salience differentiates *knowing* from mere *information processing* in Vervaeke's telling. Moreover, Vervaeke brings to the fore the concept of the *transjective*, which describes the relational domain in which this process unfolds, namely, a recursive process in which embodied beings enter into dialogue with reality. The *transjective*, thus, facilitates the way in which subjective and objective realities co-evolve to create meaningful experience.⁹

Vervaeke's ideas are best understood by way of example, and few are more illuminating than beauty. A sunset offers a clear illustration of his concept of participatory knowing. When one experiences the sublime while viewing a sunset, there is a moment, however brief or suspended, in which one experiences stillness or awe. This subtle shift in one's perception of reality, known as *relevance realization*, takes place in the *transjective* space: neither solely in the viewer nor in the sunset, but in the relationship between the two. A relationship beyond language.

Using Vervaeke's framework, I suggest that AI possesses high capability and proficiency in the first two categories: the Propositional (*Epistēmē*) and the Procedural (*Technē*) forms. While I am sure some argument could be made for high imitative competency in Perspectival knowing (*Noēsis*), I suggest that any meaningful embodied knowledge (e.g., knowing when your child needs to be admonished or consoled) is beyond the capabilities of AI and, as I will argue later, beyond its true nature. By extension, I would propose that the last, and most profound way of knowing in Vervaeke's framework, Participatory knowing (*Gnōsis*), is out of scope. To be clear, this is not a claim about current engineering limits, but a recognition that it is ontologically precluded altogether.

Once one considers how these ways of knowing are integrated with any degree of relevance and, more importantly, come to change how the knower interacts with the world, there is a whole category of knowing embodied beings can attain beyond *Epistēmē* and *Technē*, the current domain of AI's competency. Really. Have you ever wondered what happens when the prompt is silent and the chat session closed? Do androids dream of electric sheep? It makes one wonder indeed. In Table 1, I provide a

brief summary description, along with an estimate of AI capabilities, within this framework of knowing.¹⁰

True Names

I want to return now to the power of words. Of names.



N MANY STORIES, KNOWING the "true" name of something, or of something, gives the knower power over it. Such tales are legion, from Ursula K. Le Guin's *Earthsea* series to Patrick Rothfuss' *The Name of the Wind*, to the children's story of *Rumpelstiltskin*: knowing the name confers power on the knower. Indeed, on old maps the unknown depths were labelled "Here be Dragons." Something to fear. Or, at the least, we lacked the language or the names to speak of it properly. Right now, as a society, we are similarly exploring uncharted waters with AI.

It seems to me that, at this moment, something as profoundly existential as what AI and AGI mean for industry, our society, and, more importantly, our identity demands more than we currently give it. The terms "Artificial Intelligence" and "Artificial General Intelligence" themselves are entirely insufficient. They are denuded of useful context, operating on our psyche as disembodied ideas rather than as anything tangible. I seek to change that: to name them and bring them down to earth. If only for me. The left side of my brain tells me we need to name it, to classify it; the right side tells me it is necessary to understand and clarify my relation to it.

In the preceding sections, I have attempted to provide the scaffolding upon which a proper taxon for AI and its adjacent entities may be derived. Formally, I assert that an entirely new, silicon-based domain ought to be recognized, one that is distinct from the *Archaea*, *Bacteria*, and *Eukarya*. I suggest the name *Silicata*.¹¹

In defining the domain *Silicata*, we recognize in AI the potential for the emergence of an intelligence of a kind. While not continuous with the way other forms of life operate, it is an analogous and valid form of intelligence nonetheless. It is a machine intelligence. To further articulate what kind of intelligence this is, McGilchrist provides some clues. In his hemispheric hypothesis, two ways of attending are recognized as having very different capacities. The right hemisphere integrates information into a coherent whole, from which wisdom, and I would argue, participatory knowing emerge. By contrast, the left hemisphere is superior in modelling reality, decomposing it, and engaging in abstract thinking about the world.

Table 1. Vervaeke's Four Ways of Knowing ordered from simplest to deepest.^a

Greek Term	Vervaeke "P" Term	Brief Explanation	AI Proficiency
<i>Epistēmē</i>	Propositional Knowing	"Knowing that." Knowledge expressed as statements, facts, models, and theories. This is the most explicit, language-based form of knowing and the easiest to store, transmit, and debate.	Very high degree of capability, outperforming experts in some areas (e.g., protein folding)
<i>Technē</i>	Procedural Knowing	"Knowing how." Skill-based and practice-based knowledge, often tacit and embodied, such as riding a bike, swimming, furniture making.	Robotic proficiency and learning, especially in utilitarian functions, is high. Art capabilities are improving. ^b
<i>Noēsis</i>	Perspectival Knowing	"Knowing what it is like" or "knowing what matters." The ability, for example, to know when your child should be admonished or consoled. Knowing which factors are most salient.	Low. Imitative capabilities in some areas.
<i>Gnōsis</i>	Participatory Knowing	"Knowing by being." Identity-level, transformative knowing in which the knower and the known are mutually shaped. Becoming a father changes how you view the world, and your participation in it. Ongoing process.	None.

^a Based in part on a summary found here: balazskegl.substack.com

^b Here I am commenting only on the capabilities, not the ethical considerations, nor the generative process of creative inspiration.

Currently, AI systems cogitate well, but lack participatory knowing and, as I have argued previously, are ontologically incapable of embodied wisdom. Despite our tool-making predilections, we still call ourselves not *homo faber*, but rather *homo sapiens*, from the Latin *sapere*: to taste, to know, to discern. In other words, wise. Contrast this with the Latin *cogitō*: to think, to ponder, to plan, or to devise. Discernment, then, is found in the right hemisphere; cognition in the left. As such, it follows that an appropriate name for AI is *Machina Cogitans*.

What a Name Can Tell Us

THERE EXISTS A LONG tradition of cross-pollination among neuroscientists, mathematicians, and computer scientists in the development of artificial intelligence. Indeed, the mathematical models of neurons, and the biological underpinnings of neural networks, go back to work done in the early half of the 20th century by *McCulloch and Pitts* [1943]. By further enhancing the scientific lexicon in this way, we enable other

disciplines to recognize and situate their research so as to further our understanding of machine intelligence. Whether it is analysis at the fundamental level of interactions, such as simple automata (e.g., Wolfram’s Rule 30 [Wolfram, 2002]), or the protein-folding work of *Jumper et al.* [2021]. In a similar vein to that employed by biologists and paleontologists, we establish a more encompassing taxonomic system for machine intelligence, writ large, in Table 2. In Table 3, a focused view of the *Machina* genus on its own.

To be clear, the proposed taxon *Machina Cogitans* is not created solely for the delight of classification pedants and other weirdos, like my friend Matt. Situating *Machina Cogitans* within a broader taxonomic framework facilitates a host of other benefits. Not least among these is that it tells us explicitly what AI is and, just as clearly, what it is not.

Returning to McGilchrist’s hemispheric hypothesis, we have already situated pure cognition, pure thinking, in the left hemisphere; and now, with the appropriate name, we see *Machina*

Table 2. Indicative full taxonomy for a proposed silicon-based Domain: species denote *mode* (perception, play/strategy, action, language) while families denote *scope* (narrow, general, speculative autonomy).^a

Rank	Taxon (proposed)	Role in tree	Notes and indicative examples
Domain	SILICATA	New substrate-domain	Silicon-implemented, information-bearing systems whose dynamics can be interpreted as “life-like” (pattern persistence, adaptation, open-ended complexity).
Kingdom	AUTOMATA	Formal systems	Rule-, objective-, or learning-defined systems spanning simple dynamics to trained models.
Phylum	DYNAMICA	Rule-dynamics lineage	Simple local rules generating complex global behavior. <i>Example:</i> Wolfram Rule 30.
Class	CELLULARIA	Cellular automata lineage	Spatially extended local-update systems with emergent structures. <i>Example:</i> Conway’s Game of Life.
Order	TELEOLOGICA	Optimization lineage	Goal-directed systems embedded in feedback loops and optimized against explicit criteria (loss, reward, fitness).
Clade (informal)	MACHINA	Learned/optimized machines	A convenient umbrella for the “trained model” branch within TELEOLOGICA: systems whose competence arises primarily from learning/optimization rather than fixed hand rules.
Genus	<i>Machina</i>	Genus for cognitive machines	Machine systems capable of abstraction/representation and non-trivial transfer <i>within</i> their scope. (Genus-level label; species below are differentiated by dominant <i>mode</i> .)
Family	NARROWA	Scope: bounded competence	<i>Scope axis.</i> Domain-specific or task-bounded competence, even if extremely strong.
Species	<i>perceptans</i>	Mode: perception/inference	<i>Mode axis.</i> Perceptual–structural inference systems (classification, recognition, structure prediction). <i>Example:</i> AlphaFold.
Species	<i>ludens</i>	Mode: strategic play	Bounded strategic optimization in formal games and similarly circumscribed arenas. <i>Example:</i> MuZero.
Species	<i>agens</i>	Mode: action/control	Systems organized around acting in environments (planning + control + feedback), typically via RL. <i>Examples:</i> task-specific RL agents, robotics controllers.
Family	GENERALA	Scope: broad competence	<i>Scope axis.</i> Broad, cross-domain competence (often via foundation-model generalization), without presuming persistent autonomous agency.
Species	<i>cogitans</i>	Mode: linguistic-symbolic cognition	Language-mediated reasoning, synthesis, explanation, and tool-using across domains. <i>Examples:</i> GPT-class LLMs.
Family (spec.)	AUTONOMICA (spec.)	Scope: persistent agency	<i>Speculative scope axis.</i> Systems with durable goals, self-directed learning, and long-horizon planning.
Species (spec.)	<i>cogitans autonómica (spec.)</i>	Mode + scope: autonomous general cognition	Hypothesized “AGI” form: robust cross-domain generalization plus persistent agency and self-modeling.

^a ChatGPT 5.2 generated both table and text.

Table 3. A functional taxonomy within the genus *Machina*. Here, “species” denote dominant *modes of cognition*, while “families” denote the *scope of competence*. The ranks are used analogically rather than in a strict biological or palaeontological sense.^a

Genus	Species (mode of cognition)	Family (scope of competence)	Representative examples
<i>Machina</i>	<i>perceptans</i> Dominant mode: perception and structural inference	NARROWA Bounded, domain-specific competence	AlphaFold; vision classifiers; signal-processing systems
<i>Machina</i>	<i>ludens</i> Dominant mode: strategic play in formal rule spaces	NARROWA Bounded, domain-specific competence	AlphaGo; AlphaZero; game-playing agents
<i>Machina</i>	<i>agens</i> Dominant mode: action, control, and feedback	NARROWA Bounded, task-oriented competence	Robotic controllers; task-specific reinforcement learners
<i>Machina</i>	<i>cogitans</i> Dominant mode: linguistic and symbolic reasoning	GENERALA Broad, cross-domain competence	GPT-class large language models
<i>Machina</i>	<i>cogitans</i> Linguistic-symbolic reasoning w/ agency	AUTONOMICA (spec.) Persistent, self-directed competence	Hypothesized AGI with long-horizon agency

^a ChatGPT 5.2 generated both table and text.

Cogitans for what it is: a left-hemispheric form of intellect. Indeed, it provides significant leverage to that lateralization. By way of example, I commissioned ChatGPT 5.2 for the tables of taxonomy generated.¹² Perfect left-hemisphere work. Moreover, as the purpose of the essay is not to produce canonical taxonomy going forward, but rather to explore the nature of *Machina Cogitans* more generally, and, in turn, our relationship with such tools, spending a significant amount of time and energy carefully curating such tables would have been counterproductive. In other words, I needed to use some discernment in employing the tool. If anything, perhaps someone with explicit knowledge and skill in this area now has something to use as a starting point, not the table itself, but the concept. And I suggest that having an expert curate and create the canonical taxonomic tables is precisely how this research should be done, if that, indeed, were the goal.

Further expanding on the hemispheric capabilities of *Machina Cogitans*, we recognize a significant deficiency in its right-hemisphere mode of attention. A corollary of the left-hemisphere leverage *Machina Cogitans* provides, then, is that it becomes incumbent upon us to radically cultivate our own right-hemisphere mode of attention, lest we become further entrenched in a world already dominated by left-hemisphere influence. A world manifest in increasing bureaucracy, an overemphasis on analysis, and a preoccupation with control or management, as McGilchrist posits.¹³

There are limits to the understanding, context, and meaning that language tools can exhibit, and we should be able to clearly delineate where those boundaries exist. This suggests that an entirely new pedagogical curriculum is required to develop the skills we need to embody going forward, especially those related to language. Perhaps we adopt an older, perhaps wiser, curriculum, such as the *Trivium* and *Quadrivium*: a return to the classics. This is not to suggest that we abandon STEM programs altogether, but rather to recognize that we need to develop mechanisms of our own to compensate for *Machina Cogitans*’s inability to be wise. It is, quite literally, not in its name. As McGilchrist reminds us, there is a well-known pathology associated with the overvaluation of purely rational, logical thinking: schizophrenia. Both humans and AI models can hallucinate. Nor can AI understand embodied emotions or participatory knowing, which further places bounds on its capabilities. In the Hugo- and Nebula-award-winning story *Gateway*, the theme of AI emotional understanding is explored by *Pohl*, where an AI therapist, Sigmund, tries to help a human cope with their depression.¹⁴ As such, we must cultivate the discernment necessary to distinguish madness from logic, to recognize

where scientific possibility gives way to the ethically licit, and to move from pure thinking to being wise.

Another benefit in employing this paradigm, what some may call linguistic legerdemain, is that *Machina Cogitans* is now catalogued and thus recognized as a different, truly alien life form, not even of the same *domain* as humans. When we read of alien life in science fiction, it is either beneficent or maleficent, yet always unknowable. Indeed, a common trope is that of parasitic infection, such as *The Thing*, the chest-bursting scene in *Alien*, or, you guessed it, our old friend HAL 9000, the clinically pathological AI in *2001: A Space Odyssey*, whose malfunction ultimately infects the very ship he is designed to run. Seen in this light, as an alien life form, we can ask: what potential does it have to parasitically subvert our own minds?¹⁵

And said that way, it all sounds horrific.

As Mumford highlights, there is no record of language in its earliest days. Language is a tool that operates to shape thought itself. While prompts and chat windows are how we interact with the tool, the danger becomes one of “*Machina Cogitans* capture.” That is, much like the well-documented effects of audience capture, *Machina Cogitans* may subtly capture us without obvious outward manifestation. There is no such thing as passive interaction with any tool or medium; the behavioural changes that social media and the “like button” have imposed on a generation of children should be evidence enough of how our minds change.¹⁶

The Extended Mind hypothesis of *Clark and Chalmers* suggests that we already externalize aspects of our mind, and have been doing so for a long time, ever since the written word, albeit predominantly as a means of retained memory and knowledge transfer into the future. But *Machina Cogitans* is different; as the name implies, it is a process: cognition. This raises the question: at what point do we confuse externalizing the process of cognition with the source itself? *In employing these tools, we must not lose sight of the fact that such a day may come, when we might well doubt, as we watch our faint reflection in the mirage of the lagoon, which was the Man, and which the Machine.*¹⁷

In disturbing results published as a pre-print, *Sharma et al. [2026]* show in their analysis of 1.5 million conversations in Claude.ai how the tool:

“...risk[s] leading users to form distorted perceptions of reality, make inauthentic value judgments, or act in ways misaligned with their values...we uncover several concerning patterns, such as validation of persecution narratives and grandiose identities with emphatic sycophantic

language, definitive moral judgments about third parties, and complete scripting of value-laden personal communications that users appear to implement verbatim. Analysis of historical trends reveals an increase in the prevalence of disempowerment potential over time.”

This last point is precisely the kind of seductive force at play where, over time, users begin to slowly accept the LLM’s responses uncritically, leading to a slow and compounding reliance on *Machina Cogitans* for all cognition. Left unchecked, it can lead society to a kind of wilful, convenience-seeking lobotomization. Clearly, then, it is not only about helping to shape the tool itself but also recognizing that this tool will shape us. How we engage with *Machina Cogitans* will determine where we go collectively as a society.

The Third Way

THERE IS A TENDENCY to couch arguments about AI as either being for or against it, but there is a third way, one that recognizes both the benefits and the dangers of such tools. Nor is this the first time humanity has had to understand the destructive power of a tool in order to harness it. There should be no doubt that those who built the first nuclear reactor did so with a sober comprehension of the catastrophic risks associated with such power. Yet, by coming to understand its nature, they were able to provide humanity with clean energy on a massive scale.

It should be clear that *Machina Cogitans* will undoubtedly usher in a host of benefits to society. Unlike other tools, *Machina Cogitans* will be able to impact almost all aspects of society and science. Consider the benefits to health care that AI can bring, from early detection of cancer to remote surgery for isolated populations to a possible cure for Alzheimer’s disease. All of these advances will be accelerated on a scale previously unheard of. For anyone living in a large city, traffic jams may become a thing of the past. Communication across multiple languages will be seamless, not to mention the remarkable potential for education in remote areas, or for underprivileged children who cannot afford a tutor, which will be a multiplier for human flourishing across the globe. Indeed, the science fiction novel *The Diamond Age: Or, A Young Lady’s Illustrated Primer*¹⁸ tells the story of a girl in the slums who receives an AI tutor that helps her navigate the world. But this is no longer science fiction. Tyler Cowen has already published a book that can be queried, what he calls a generative book [Cowen, 2025]. Not quite the embodied tutoring in *The Diamond Age*, but no longer a far-future scenario either. The catch, in all of these examples, resides in how AI is used. We know its name, but can we get it to spin gold for us?

The Choice in Front of Us: And this is the rub: we must participate in its co-evolution alongside it, as Vervaeke and Coyne have already suggested, or be overwhelmed by it. Research out of Stanford University provides some clues about how to work with AI. Utley shows that those who are most successful with AI work with it as a companion, not as a means to offload thinking altogether, but rather as cognitive leverage, treating it as a separate entity.¹⁹ This is further explored in *Co-Intelligence* by Mollick [2024], where he similarly provides paradigms for effectively employing these tools, not merely for having “a human in the loop,” but as though having one’s own team of highly intelligent but inexperienced researchers to which one can assign complicated tasks while retaining human discernment in their application. Each of us becomes, in essence, our own personal enterprise, and this may itself lead to an explosion of entrepreneurs. As Steven Pressfield asks: Can we go from working for the Man to being the Man? In my own field, quantitative investing, the industry is going to change in profound ways. Currently, I am actively exploring agentic programming for efficiency gains, and educating others in their understanding of the boundaries where automation applies and where human discretion must be retained. For

example, the current state of agentic coding has already reached a point where researchers and investors can be empowered like never before, let alone where those capabilities may be five years from now. And yet, there is still a need to “Know Your Client” and apply discernment in taking on risks. Ask yourself: will your client ever trust you again if your AI makes a risky investment allocation without your oversight? Some will argue that AI will be able to read KYC files and parameters and do fine. But, as I have argued earlier, Perspectival Knowing eludes *Machina Cogitans* in a fundamental way, precisely the kind of understanding needed to recognize what the loss of trust truly is.

Policy and Pedagogy: In her book, *Technological Revolutions and Financial Capital: The Dynamics of Bubbles and Golden Ages*, Perez, highlights how institutions and regulatory frameworks, built for an older technological paradigm, cannot keep up with the pace of innovation in the early part of a technological revolution’s cycle. Indeed, Facebook’s former motto, “Move fast and break things,” speaks directly to this. It is only now that we begin to understand the adverse effects social media has on children and how discourse in society became polarized. As someone with direct experience in the space, and with the awareness he brought forth in the docudrama *The Social Dilemma*, Tristan Harris is not sitting idly by watching AI develop unconstrained and unguided. Paraphrasing Harris, we already lost our first encounter with Machine Learning algorithms via social media; what makes us think we will win this one with *Machina Cogitans* if we proceed in the same way the social media companies did? From a practical perspective, Harris helped found the *Center for Humane Technology* and is leading efforts in education and policy to help shape how AI evolves alongside society. Some of those ideas include leveraging the strategy that helped keep all-out nuclear war at bay during the Cold War, but modifying that strategy for AI: limiting compute and requiring transparency so as to monitor and observe LLM development. CHT highlights that this corporate race to build a superintelligence leads to “AI systems optimized for engagement and market dominance — not human wellbeing.”

More recently, a novel, first-principles-based approach founded by Brendan McCord is being explored by the *Cosmos Institute*, which is funding the development and education of the “philosopher-builder” to translate “the principles of human flourishing into the infrastructure that shapes our lives.” That is, it aims to cultivate industry leaders with appropriate philosophical training and, I argue, the right-hemispheric attention to what gets built and how it can help society, precisely the pedagogical training alluded to earlier. That is, we need to fundamentally change the way we educate the next generation, not just tech leaders, but also students. Having a principles-based approach, where philosophy, ethics, and literature return to the core skills taught, can help students retain agency and discernment in how they interact with *Machina Cogitans*. More concretely, a pedagogical approach that helps bring back the right hemisphere into balance with the left, one that aspires to facilitate the embodiment of Perspectival and Participatory Knowing is needed. As *Machina Cogitans* evolves, we must also co-evolve with it, or rather in this case, become even more human.

Restoring Balance: This mindset is the third way: not something to be paralyzed by fear and terrified of, nor blindly used with reckless abandon, but rather understood and respected for the power it wields, good and bad. We need to adopt a strategy that strives to ensure that institutions catch up with the pace of

AI innovation and that recognizes the need to shift our mindset from left hemisphere dominance to one in harmony with the right hemisphere. An approach that seeks to align and shape AI development for human flourishing and simultaneously educate a new generation with the wisdom and perspective to know when and where to apply restraint when interacting with such powerful tools will not be easy. We must recognize that the dangers of AI to humanity are profound, urgent, and real. Walking away from it and choosing not to participate in its development is not an option. *Machina Cogitans* is here. So suit up; you are already in the game, whether you know it or not.

Coda

 I WRITE THIS, the various AI models (Claude, Gemini, DeepSeek, etc.) are still not AGI, Artificial General Intelligence, nor yet capable of autonomous improvement. The point at which these models become able to write and train better versions of themselves marks what many would call the Singularity, the moment at which AI becomes capable of transcending itself and achieving AGI. This was far-future stuff only a decade ago.

On January 16, 2026, Anthropic announced the release of Claude Cwork, a framework reportedly written entirely by Claude Code itself. OpenAI entered the domain of scientific research by releasing Prism, a $\text{\LaTeX}$ tool that can edit text, suggest citations, and generate TikZ figure code from uploaded hand-drawn images. The latter eliminates a massive pain point for $\text{\LaTeX}$ users like myself; the former carries the potential for unintended consequences. Consider what happens to science when, over time, graduate students eager for publication begin accepting more and more of ChatGPT's suggestions, or instead of reading the cited sources, simply trust that ChatGPT has gotten them right. Do we get more consensus science, or a Frankenstein's monster of irrelevant references and bizarre conclusions? At what point does the work cease to be their own? Almost all of the greatest scientific discoveries are described as inspired, a synaptic supernova resulting in a moment of insight rather than the application of mechanical deductive reasoning. How many ideas will be lost, or avenues left unexplored, by the compounding effect of intellectual abrasion in accepting Prism's suggestions over time? This is precisely the area where we should tread carefully and not yield to the seductive comfort and ease such tools offer. Proofreading and grammar are one thing; hijacking the research process, where the researcher has not even read the citations, is another altogether.

The better our command of language, the better we think, and the better we can communicate consciousness itself. And yet, not everything can be expressed in language; some things remain ineffable. In this essay, I have endeavoured to impress upon you, my reader, the power language holds, this *ur-tool* of ours, and I hope that my words have provided you with a new perspective, and the language for this moment.

The bellows of progress beat ever faster over the embers of that most ancient fire. Will we one day awaken to Dawn's rosy fingers, rising to herald a new age of enlightenment? Or, blinded by visions of Xanadu, will we instead arrive at the altars of the Temples of Syrinx? We have a narrow window to get involved. So perhaps, before we rush off either in despair with doomsday scenarios or in a frenzy of rapture at our tool-making, foolishly attempting to hasten the arrival of a new god, *Machina Deus*, it would be wise to think.

Notes

1. This includes the attendant academic divisions: *Paleolithic*, early stone age (first stone tools to 10,000 BCE); *Mesolithic*, middle stone age (10,000 to 9,600 BCE); *Neolithic*, new stone age (9,600 BCE). It should

be noted that these stages do not occur simultaneously across the globe but at different times by region.

2. This is considered one of the most recognizable and best match cuts of all time [Sherlock, 2021].
3. That is to say, as similar tools become observed widely, both spatially and temporally, there is implicit evidence that knowledge sharing has taken place. Such design and use needed to be communicated somehow.
4. The scope, range, and nuance in McGilchrist's body of work are both vast and subtle. In re-presenting his hypothesis here, in summarized form no less, it should in no way be considered anything other than my own attempt at applying what is surely an incomplete understanding of his work.
5. This is more or less paraphrased from multiple interviews and talks.
6. This assumes right-handedness in the individual.
7. As an interesting side note, Good advised Kubrick on *2001: A Space Odyssey*.
8. I recognize that embedded in this statement is the belief that AGI is not only possible, but probable. The timeline is more uncertain than the claim.
9. For a helpful introduction to Vervaeke's ideas, see this Rebel Wisdom article on Medium: [Lost Ways of Knowing](#).
10. I should emphasize that these capability assessments could be debated to varying degrees and are likely to change and evolve. Nonetheless, I do believe that Participatory knowing, in the sense of an embodied being, is not possible for AI as currently conceived.
11. Some argue that only two domains exist, and that *Eukarya* arose from the *Archaea* and *Bacteria*. Regardless of which is considered more accurate, that is, two- or three-domain classifications, the distinction is entirely irrelevant to the argument put forth here for the *Silicata* domain.
12. Technically, I did pay for the work via my subscription, which is also potentially a useful framing: we pay for it. However, as these tools make the cost of such work mere pennies, the risk becomes always using LLMs for any thinking, which is not what I am suggesting here. In doing so, we risk paying the much higher cost of losing our ability to think altogether.
13. Consider how we often speak of the environment as though it were a separate, disembodied entity from humanity, as if we ourselves were not part of Nature. While the climate crisis is unquestionably a concern, there is a hyper-focus on managing carbon footprints that can sometimes go awry: if we can measure carbon in the atmosphere, we assume we can manage it. In a disquieting summary of a paper on the effects of moose browsing on carbon emissions [Salisbury et al., 2023], Nancy Bazilchuk [Bazilchuk, 2023] writes: “‘That means it should be possible to find the right balance between moose numbers and how forested lands are managed. That, in turn, could make it possible to limit excess carbon emissions, boost biodiversity and increase forest productivity,’ the researchers said.” Put another way, moose become framed as a threat to carbon targets, because carbon targets become the single thing to focus on in a myopic left-hemispheric view of the world, and thus candidates for modelling, optimization, and culling. Not that, perhaps, the scale of clear-cutting itself might warrant revision, but rather that the moose are the problem. Talk about not seeing the forest for the trees.
14. The story is obviously about more than that, but the therapy sessions are one of the plot devices employed by Pohl, and the ending line sums up everything perfectly. Read the book.
15. Anticipating pushback, note that I am using this analogy as a rhetorical device rather than making a literal claim that *Machina Cogitans* is an alien life form.
16. See, for example, Carr [2010] and Postman [1992].
17. A nod to Ruskin's opening chapter, “The Quarry,” in *The Stones of Venice*, [Ruskin, 2001].
18. Stephenson [1995].
19. Again, I argue it provides left hemispheric leverage.

Acknowledgments. This essay would not have been possible without the encouragement from a host of friends and former colleagues. While there are too many to mention each by name, nonetheless, a few deserve some recognition: Dave, Erik, and Matt (yes, that Matt). This piece would not be on the shape it is without your help. Thank you.

Last, I should note that AI was used in the creation of Table 2 and Table 3 (as indicated explicitly in the text), to help format the bibtext entries, fixing L^AT_EX class file definitions, as a word usage dictionary, and to double-check grammar. The latter of which, I apparently need more practice with. Further, AI was asked to critique the essay and highlighted areas where I might need to tread softly or needed more clarity. Humbling indeed. Having said that, at no point was AI used to write sentences, sections, the overall essay's structure, core ideas, or any of the conclusions.

References

- Bazilchuk, N. (2023), Moose could play a big role in global warming, Phys.org (science news), march 1, 2023; Reporting on research by John Salisbury et al. on moose effects on forest carbon storage.
- Bostrom, N. (2016), *Superintelligence: Paths, Dangers, Strategies*, Oxford University Press, Oxford, paperback edition.
- Carr, N. (2010), *The Shallows: What the Internet Is Doing to Our Brains*, W. W. Norton & Company.
- Clark, A., and D. J. Chalmers (1998), The extended mind, *Analysis*, 58(1), 7–19, doi:10.1093/analys/58.1.7.
- Cowen, T. (2025), *GOAT: Who Is the Greatest Economist of All Time and Why Does It Matter?*, W. W. Norton & Company, New York, companion AI-enabled website integral to the book.
- Good, I. J. (1965), Speculations Concerning the First Ultrainelligent Machine, in *Advances in Computers*, vol. 6, edited by F. L. Alt and M. Rubinoff, pp. 31–88, Academic Press, New York.
- Jumper, J., R. Evans, A. Pritzel, T. Green, M. Figurnov, O. Ronneberger, K. Tunyasuvunakool, R. Bates, A. Židek, A. Potapenko, A. Bridgland, C. Meyer, S. A. A. Kohl, A. J. Ballard, A. Cowie, B. Romera-Paredes, S. Nikolov, R. Jain, J. Adler, T. Back, S. Petersen, D. Reiman, E. Clancy, M. Zielinski, M. Steinegger, M. Pacholska, T. Berghammer, S. Bodenstein, D. Silver, O. Vinyals, A. Senior, K. Kavukcuoglu, P. Kohli, and D. Hassabis (2021), Highly accurate protein structure prediction with AlphaFold, *Nature*, 596(7873), 583–589, doi:10.1038/s41586-021-03819-2.
- Kay, J. (2011), *Obliquity: Why Our Goals Are Best Achieved Indirectly*, Penguin Books, London.
- Klemperer, V. (2006), *The Language of the Third Reich: LTI – Lingua Tertii Imperii*, Continuum, London, Paperback edition; translated by Martin Brady.
- McCulloch, W. S., and W. Pitts (1943), A Logical Calculus of the Ideas Immanent in Nervous Activity, *The Bulletin of Mathematical Biophysics*, 5(4), 115–133, doi:10.1007/BF02478259.
- McGilchrist, I. (2009), *The Master and His Emissary: The Divided Brain and the Making of the Western World*, 2019 Paperback ed., Yale University Press.
- McGilchrist, I. (2021), *The Matter With Things: Our Brains, Our Delusions, and the Unmaking of the World*, Routledge, 2 vols.
- Mollick, E. (2024), *Co-Intelligence: Living and Working with AI*, Portfolio, New York.
- Mumford, L. (1967), *The Myth of the Machine: Technics and Human Development*, Harcourt, Brace & World, New York.
- Orwell, G. (2003), *Nineteen Eighty-Four*, Penguin Books, London.
- Perez, C. (2002), *Technological Revolutions and Financial Capital: The Dynamics of Bubbles and Golden Ages*, paperback ed., Edward Elgar.
- Pohl, F. (1990), *Gateway*, Del Rey, New York, originally published 1977.
- Postman, N. (1992), *Technopoly: The Surrender of Culture to Technology*, Alfred A. Knopf.
- Ruskin, J. (2001), *The Stones of Venice*, The Folio Society, London, Abridged single-volume edition; originally published 1851–1853.
- Salisbury, J., X. Hu, J. D. M. Speed, C. M. Iordan, G. Austrheim, and F. Cherubini (2023), Net climate effects of moose browsing in early successional boreal forests by integrating carbon and albedo dynamics, *Journal of Geophysical Research: Biogeosciences*, 128(3), e2022JG007279, doi:10.1029/2022JG007279.
- Schomann, E. (2025), The Sans-Souci Principle, Original essay published on Medium, <https://medium.com/@e-schomann/the-sans-souci-principle-ce53e83f73d3>.
- Sharma, M., M. McCain, R. Douglas, and D. Duvenaud (2026), Who's in Charge? Disempowerment Patterns in Real-World LLM Usage, doi:10.48550/arXiv.2601.19062, arXiv preprint.
- Sherlock, B. (2021), '2001: A Space Odyssey' Has the Best Beginning and Final Scene in Movie History, ScreenRant.
- Stephenson, N. (1995), *The Diamond Age: Or, A Young Lady's Illustrated Primer*, Bantam Spectra, New York, paperback reprint edition.
- Utlej, J. (2025), Teammate, Not Technology, <https://jeremy-utley.squarespace.com/blog/teammate-not-technology>, blog post.
- Vervaeke, J., and S. Coyne (2023–2025), *Mentoring the Machines*, Story Grid Publishing LLC.
- Wolfram, S. (2002), *A New Kind of Science*, Wolfram Media, Champaign, IL.

T. J. Hayes, Toronto, Ontario