

IOCPH
2025
Conference

The 1st International Online Conference of the Journal Philosophies

10–14 June 2025 | Online



Program and Abstract Book

Organizers



philosophies

Media Partners



entropy



information



logics



Organizing Committee

Conference Chairs

Prof. Dr. Marcin J. Schroeder
Prof. Dr. Gordana Dodig Crnkovic

Keynote Speaker

Dr. Stephen Wolfram

Invited Speakers

Assoc. Prof. Marcin Milkowski
Dr. David Gamez
Dr. Hector Zenil
Prof. Dr. Andrew Adamatzky
Prof. Dr. Diane Proudfoot
Prof. Dr. Jordi Vallverdú
Prof. Dr. Lorenzo Magnani
Prof. Dr. Michael Levin
Prof. Dr. Oron Shagrir
Prof. Dr. Selmer Bringsjord
Prof. Dr. Vincent C. Müller

Organised by



Conference Secretariat Members

Ms. Allie Shi
Ms. Coco Yang
Ms. Riley Liu
Email: iocph2025@mdpi.com



10–14 June 2025
Online

The 1st International Online
Conference of the Journal Philosophies

Table of Contents

Welcome from the Chair	1
Keynote Speaker	7
Invited Speakers	8
Program Overview	12
IOCPH 2025 Program	13
General Session	19

Welcome from the Chair

In an era of unprecedented technological advancements in intelligent and cognitive technologies, the concept of intelligence—both human and artificial—takes center stage. **The journal *Philosophies* is proud to announce its upcoming conference, “Intelligent Inquiry into Intelligence.”** This event is designed to transcend the superficial discourse often found in popular media and focus heavily on the essence and implications of intelligence in its multiple contexts transcending the most common comparisons of human and technological capacities. Some of these contexts (such as artificial intelligence or advanced alien intelligence) generate popular interest among the general audience, while others are familiar to only a few groups of experts (e.g., intelligence characterizing multiple forms of life starting from a cellular level). The conference aims to set conceptual and methodological foundations for productive interactions and cooperation between diverse directions of inquiry and to foster a thoughtful and rigorous examination of intelligence, supported by interdisciplinary research and philosophical reflection.

The term “intelligence” has become ubiquitous, often overshadowed by buzzwords and misconceptions. Robert Louis Stevenson wrote in 1881, “Man is a creature who lives not upon bread alone, but principally by catchwords [...]” [1] (41). The same can be said today, possibly with the new expression “buzzwords” instead of “catchwords”.

This conference aims to cut through the confusion and myths that currently surround AI and intelligence in general, promoting informed discussions that can shape future policies, educational frameworks, and societal impacts. By bringing together leading philosophers and researchers from diverse fields, we seek to build a solid foundation for the study of intelligence that moves beyond the hype and toward profound insights, striving to achieve a deeper, more nuanced understanding of intelligence in our rapidly changing world.

Following the mission of *Philosophies*, the general goal of this conference is to promote interaction and mutual stimulation between philosophical reflection, scientific research, and other forms of intellectual inquiry, including diverse forms of expression for philosophical reflection. The intended subject of these diverse inquiries is intelligence, but in the absence of a commonly accepted definition for this omnipresent but elusive concept, intelligence can be considered within the context of the following configuration of themes: information, knowledge, rationality, logic, computation, complexity, creativity, autonomy, agency, life, cognition, and consciousness. None of these ideas have achieved the status of a clearly, uniformly defined, and well-understood concept that can be used as a stepping stone for the study of the others. Instead, they constitute a family of concepts that is a mutually interdependent network. Even information that is typically considered more general than the other ideas on the list seems to fit the role of a genus for definitions well, and has multiple competing conceptualizations and unsettled attempts to equip it with semantic characteristics.

A secondary but closely related goal supporting the integration of diverse forms of inquiry of intelligence is to prevent confusion created by the diverse contexts of these inquiries. It is a popular but dangerous assumption that in the absence of definitions of concepts that escape rigorous analysis, we can use their common-sense understanding. The danger is in the illusion of commonality. Empirical studies have shown that common sense turns out to not be very common (not in the Voltairian sense of its rarity, but as the lack of uniform understanding) [2]. Thus, the claim that common-sense communication secures mutual understanding may be an

old illusion already known to Socrates: the use of the same words does not imply the same meaning. Here, philosophy helps in finding, if not commonly accepted definitions, interpretive methods of hermeneutics.

We would like to bring to the attention of all participants the need to create a common ground that incorporates **intelligence (artificial and natural, including human) and its entire network of concepts including information, knowledge, rationality, logic, computation, complexity, creativity, autonomy, agency, life, cognition, awareness, and consciousness**. This means that whenever these concepts are employed in the study of intelligence, it is necessary to specify their understanding without the assumption that their meaning is self-explanatory. The conference aims to bring together diverse perspectives on intelligence, and this means that its relations with other relevant concepts are of special importance. All studies of the configuration of concepts will contribute to this goal.

To facilitate mutual interactions between different forms of inquiry of intelligence, contributors should avoid the use of specialistic jargon and the works presented at the conference involving mathematical, statistical, or other scientific formalisms should be accompanied by sufficiently extensive explanations and interpretations to be comprehensible to the wider audience. The expectation that such an explanation is possible follows from the famous claim by Richard Feynmann in the context of advanced theoretical physics that if you cannot explain something in a way comprehensible to a freshman student, you most likely do not understand it yourself [3].

The following outstanding questions of the Intelligent Inquiry into Intelligence are stated here as examples, rather than limits, of interest for submissions to the conference. The references are provided **not necessarily as patterns to follow** in submissions but rather as sources of information about the content of concisely formulated questions.

Questions about

- **The epistemological status of intelligence** and the systems qualified as intelligent.
- **The axiological issues of intelligence** (value, in general, rather than moral values in particular, emphasizing the plurality and heterogeneity of values) and the systems qualified as intelligent.
- **The ontological status of intelligence** and the systems qualified as intelligent.
- The related issue of **the status of collective or distributed intelligence** as opposed to the traditional view of intelligence as an exclusive characteristic of singular entities with central architecture [4-6].
- **The role of logic (or possibly alternative logics, rationality) in intelligence.**
- **Non-human natural intelligence** (the wide spectrum of views from the claim cognition = life to the various forms of plant, fungi, and animal cognition, to the denial of non-human natural intelligence or the claim of an essential qualitative difference between human and non-human intelligence) [7, 13].
- **Non-conventional computing** and its role in understanding intelligence and the design of intelligent systems [14-16].
- **The function of non-conventional computing by non-human animals and other life forms** in their existence and its qualification as a form of intelligence. [14-16].
- The meaning of the **separate qualification of natural and artificial intelligence** (including the paradox of natural human intelligence as opposed to artificial intelligence created by humans with their natural status) [16].

- The relationship between **intelligence, agency, and autonomy (e.g., their mutual roles as necessary conditions for others)**.
- **Intelligence as normative vs. descriptive attribution** (including the paradox of the contradiction between the attribution of intelligence as a distinctive characteristic of some human individuals and the claim that the main difficulty in the design of authentic artificial intelligent systems is to equip them with common sense) [2,17–19].
- **The capacity of common sense for/in intelligent systems** sought in technological research (a diverse range of concepts from the Aristotelian capacity of humans and animals to coordinate senses, through a variety of expressions in different languages to the philosophical positions on common sense of Wittgenstein and Turing, to the concept of AGI) [2,17–19].
- Involvement of **causality and determinism** in the conceptualization of intelligence (e.g., causality as opposed to interactionism, or need for the transition from Generative to Causal AI). [20].
- The relationship between **intelligence and creativity** (including tests such as the Lovelace test in the context of artificial intelligence) [21–24].
- The relationship between **intelligence and consciousness**.
- The related issue of **non-human natural and artificial intelligent systems equipped with consciousness** (including the recently published The New York Declaration on Animal Consciousness) [25].
- Diverse theories and **models of consciousness** (from panpsychism of IIT to more general models of information integration, to alternative approaches) and their relation to intelligence [26–35].
- Diverse understanding of **intelligence in multicultural, cross-cultural, and inter-cultural perspectives** [36–40].
- Learning about understanding non-human intelligence from the diverse cultural understanding of human intelligence. [36–40].
- **Emic and etic methodologies** of the multi-cultural studies of intelligence [36–40].
- **Nature and/or nurture of intelligence** (including the possibility of growing intelligence through learning).
- **The technology of intelligent systems** and its impact on individuals, collectives, society, and humanity (including alternative technological architectures and platforms).
- **Epistemological tools for the study of intelligence** and use of **the technology of intelligent systems as an epistemological tool** for the exploration of reality [41–43].
- The need for **transformation/evolution of human ethics to reflect the role of technology (in particular intelligence-related software and hardware)** in human affairs and its influence on the human condition. [44].
- The relationship between the study of intelligence (natural or artificial) and Contemporary Natural Philosophy. [40, 45–47].

The final, but possibly most important, is **the question about possible non-anthropomorphic intelligent systems**. Artificial intelligence and even natural non-human intelligence are usually (almost always) considered from the perspective of human intelligence with its multiple variations derived from human experience, values, and ideas. All searches for extra-terrestrial intelligence have failed thus far. Is it possible that the failure is a result of looking for the wrong manifestations of non-human intelligence?

The list of philosophical questions is not exhaustive. Contributions with other philosophical perspectives are welcome but considering the profile of the journal *Philosophies* sponsoring the conference, it is expected that they all serve the goal of enriching philosophical, intelligent inquiry into intelligence.

References

1. Stevenson, R. L. *Virginibus Puerisque and Other Papers*. Kegan Paul, London, 1881. Retrieved on April 23, 2024 from *Virginibus puerisque*, and other papers : Stevenson, Robert Louis, 1850-1894 : Free Download, Borrow, and Streaming : Internet Archive.
2. Whiting, M. E. and Watts, D. J. A Framework for Quantifying Individual and Collective Common Sense. *PNAS*, 2024, 121(4), e2309535121, <https://doi.org/10.1073/pnas.2309535121>
3. Goodstein, D. L. Richard P. Feynmann, Teacher. *Physics Today*, **1989**, 42(2), 70. <http://dx.doi.org/10.1063/1.881195>
4. Levin, M. (2023). Collective Intelligence of Morphogenesis as a Teleonomic Process, eds. Corning, P. A., Kauffman, S. A., Noble, D., Shapiro, J. A., Vane-Wright, R. I., and Pross, A. *Evolution "On Purpose": Teleonomy in Living Systems*, MIT Press, Cambridge, MA, USA, pp. 175-197.
5. McMillen, P., Levin, M. Collective intelligence: A unifying concept for integrating biology across scales and substrates. *Commun Biol* **2024**, 7, 378. <https://doi.org/10.1038/s42003-024-06037-4>
6. Friston, K. J. Designing Ecosystems of Intelligence From First Principles. *Collective Intelligence* **2024**, 3(1), 1-18.
7. Stewart, J. Cognition = Life: Implications for Higher-Level Cognition. *Behavioural Processes* **1996**, 35, 311-326.
8. Levin, M. Technological Approach to Mind Everywhere (TAME): An Experimentally Grounded Framework for Understanding Diverse Bodies and Minds. arXiv:2201.10346v1 [q-bio.TO] <https://doi.org/10.48550/arXiv.2201.10346>
9. Adamatzky, A. Towards fungal computer. *Interface Focus*, **2018**, 8(6), 20180029.
10. Adamatzky, A. Plant leaf computing. *Biosystems*, **2019**, 182, 59-64.
11. Adamatzky, A. Towards fungal computer. *Interface focus*, **2018**, 8(6), 20180029.
12. Vallverdú, J., Castro, O., Mayne, R., Talanov, M., Levin, M., Baluška, F., Gunji, Y.-P., Dussutour, A., Zenil, H. & Adamatzky, A. Slime mold: the fundamental mechanisms of biological cognition. *Biosystems* **2018**, 165, 57-70.
13. Vallverdú, J. What if plants compute? Commentary on Sgundo-Ortin and Calvo 'Plant Sentience'. *Animal Sentience* **2023**.471
14. de Silveira, T. B. N. and H. S. Lopes, H. S. (2023). Intelligence across humans and machines: a joint perspective, *Front. Psychol.*, 14: 120761. doi10.3389/fpsyg.2023.1209761.
15. Levin, M. Technological Approach to Mind Everywhere (TAME): An Experimentally Grounded Framework for Understanding Diverse Bodies and Minds. arXiv:2201.10346v1 [q-bio.TO] <https://doi.org/10.48550/arXiv.2201.10346>
16. Holm, S. & Powell, R. Organism, machine, artifact: The conceptual and normative challenges of synthetic biology. *Studies in History and Philosophy of Biological and Biomedical Sciences* (**2013**), <http://dx.doi.org/10.1016/j.shpsc.2013.05.009>
17. Davis, E. and Marcus, G. Commonsense Reasoning and Commonsense Knowledge in Artificial Intelligence. *Communications of the ACM*, September **2015**, 58(9), 92-103. DOI:10.1145/2701413.
18. Floyd, J. Turing on "Common Sense": Cambridge Resonances. In: Floyd, J., Bokulich, A. (eds) *Philosophical Explorations of the Legacy of Alan Turing*. Boston Studies in the Philosophy and History of Science, vol 324. Springer, Cham. **2017**, pp. 52. https://doi.org/10.1007/978-3-319-53280-6_5
19. Harris, R. Integrationism, Language, Mind, and World. *Language Sciences* **2004**, 26, 727-739.
20. Sgaier, S. K., Huang, and Grace, C. The Case of Causal AI. *Stanford Social Innovation Review*, **2020**, 18(3), 50-55.
21. Boden, M. A. Creativity and Artificial Intelligence. *Artificial Intelligence* **1998**, 103, 347-356
22. Bringsjord, S., Bello, P., and Ferrucci, D. Creativity, the Turing Test, and the (Better) Lovelace Test. *Mind and Machines* **2001**, 11, 3-27.
23. Zhou, E., and Lee, D. Generative Artificial Intelligence, Human Creativity, and Art. *PNAS Nexus*, **2024**, 3(3), 1-8.
24. Editorial. Collaborative Creativity in AI. *Nature Machine Intelligence* September **2022**, 4, 733.
25. The New York Declaration on Animal Consciousness. The New York Declaration on Animal Consciousness (google.com)
26. Tononi, G. The information integration theory of consciousness, eds. Velmans, M. and Schneider, S. *The Blackwell Companion to Consciousness*, Blackwell, Malden, MA, USA, **2007**, pp. 287-299.
27. Tononi, G., Sporns, O., and Edelman, G. M. A measure for brain complexity: Relating functional segregation and integration in the nervous system. *Proc. Natl. Acad. Sci. USA*, **1994**, 91, pp. 5033-5037.
28. Tononi, G. and Edelman, G. M. Consciousness and complexity. *Science*, **1998**, 282, pp. 1846-1851.
29. Reber, A. S. The CBC Theory and Its Entailments. *EMBO reports*, January **2024**, 25, 8-12.

30. Fleming, S. *et al.*, The Integrated Information Theory of Consciousness as Pseudoscience. Preprint at PsyArXiv, **2023**, <https://doi.org/10.31234/osf.io/zsr78>.
31. Lenharo, M. Consciousness theory slammed as pseudoscience - sparking uproar: Researchers publicly call out theory that they say is not well supported by science, but that gets undue attention. *Nature NEWS* 20 September **2023**, available online at Consciousness theory slammed as 'pseudoscience' – sparking uproar (nature.com)
32. Cogitate Consortium *et. al.* An adversarial collaboration to critically evaluate theories of consciousness. bioRxiv **2023**.06.23.546249; doi: <https://doi.org/10.1101/2023.06.23.546249>
33. Lau, H. (2023). What is a Pseudoscience of Consciousness? Lessons from Recent Adversarial Collaborations. September 17, **2023**. <https://doi.org/10.31234/osf.io/28z3y>
34. Lenharo, M. AI Consciousness: Scientists Say We Need Answers - Researchers urge more funding to study the boundary between conscious and unconscious systems. *Nature*, 11 January **2024**, 625, 226.
35. Bayne, T. *et al.* Tests for Consciousness in Humans and Beyond. *Trends in Cognitive Sciences* **2024**, 2533 (in print)
36. Markus, H.R.; Kitayama, S. Culture and the self: Implications for cognition, emotion, and motivation. *Psychol. Rev.* **1991**, *98*, 224–253.
37. Frake, C. O. Cultural Ecology and Ethnography. *American Anthropologist*, **1962**, *64*, 53–59.
38. Frake, C. O., The Ethnographic Study of Cognitive Systems. In Gladwin, T. & Sturtevant, T. W. C. (Eds.) *Anthropology and Human Behavior*, Anthropological Society of Washington, Washington D.C., **1962**.
39. Goodenough, W. H. Componential analysis and the study of meaning. *Language*, **1956**, *32*(1), 195–216.
40. Simeonov, P. L., Gare, A., Matsuno, K. & Igamberdiev, A. U. Editorial. Special issue on Integral Biomathics: The Necessary Conjunction of the Western and Eastern Thought Traditions for Exploring the Nature of Mind and Life. *Progress in Biophysics and Molecular Biology* December **2017**, *131*, Focussed Issue):1–11.
41. Hoom, J. F. & Chen, J. J-Y. Epistemic considerations when AI answers questions for us. arXiv:2304.1452v1 [cs.CY] **2023**. <https://doi.org/10.48550/arXiv:2304.1452v1>
42. Wheeler, G. R. & Moniz Pereira, L. Epistemology and artificial intelligence. *Journal of Applied Logic* 2004, *2*(4), 489–493.
43. Russo, F., Schliesser, E., & Wagemans, J. Connecting ethics and epistemology of AI. *AI & Soc.* **2023**, <https://doi.org/10.1007/s00146-022-01617-6>
44. Bankins, S. & Formosa, P. The Ethical Implications of Artificial Intelligence (AI) For Meaningful Work. *J. Bus. Ethics* **2023**, *185*, 725–740.
45. Kauffman, S. A. & Roli, A. A Third Transition in Science? **2023** royalsocietypublishing.org/journal/rsfs *Interface Focus* *13*: 20220063
46. Kauffman, S. Is There a Fourth Law for Non-Ergodic Systems That Do Work to Construct Their Expanding Phase Space? *Entropy* **2022**, *24*, 1383. <https://doi.org/10.3390/e24101383>
47. Simeonov, P. L. *et al.* Stepping Beyond the Newtonian Paradigm in Biology: Towards an Integrable Model of Life – Accelerating Discovery in the Biological Foundations of Science INBIOSA White Paper. In P. L. Simeonov, L. S. Smith, A. C. Ehresmann (Eds.) *Integral Biomathics: Tracing the Road to Reality*. Springer, Berlin, **2012**, pp. 319–418.



Prof. Dr. Marcin J. Schroeder
Conference Chair

1. Professor Emeritus, Akita International University, Akita, Japan,
2. Editor-in-Chief, the journal *Philosophies of MDPI*



Prof. Dr. Gordana Dodig Crnkovic
Conference Chair

1. Professor of Computer Science, Mälardalen University, Västerås, Sweden
2. Professor of Interaction Design, Chalmers University of Technology, Gothenburg, Sweden



philosophies

Philosophies is an international, peer-reviewed, open access journal promoting re-integration of diverse forms of philosophical reflection and scientific research on fundamental issues in science, technology and culture, published bimonthly online by MDPI. The International Society for Information Studies (IS4SI) is affiliated with *Philosophies* and their members receive a discount on the article processing charge.

Journal Webpage: <https://www.mdpi.com/journal/philosophies>

Keynote Speaker



Dr. Stephen Wolfram

Founder & CEO of Wolfram Research,
Creator of Mathematica,
Wolfram|Alpha & Wolfram Language,
Author of A New Kind of Science and
other books,
Originator of Wolfram Physics Project,
<https://www.stephenwolfram.com/about/>

Computational Foundations of Minds and the Universe

Wolfram's recent work on the foundations of physics, mathematics, biology and machine learning introduces a major new framework for thinking about fundamental philosophical questions. This talk will provide a non-technical, philosophically oriented survey of these directions.
<https://writings.stephenwolfram.com/category/philosophy/>

Invited Speakers



**Prof. Dr. Andrew
Adamatzky**

University of the West
of England, Bristol, UK

Origin of Intelligence

What is intelligence, and how deeply is it rooted in the fabric of life itself? In this talk, I explore the fundamental emergence of intelligent behaviours far beyond the animal kingdom. Beginning with spiking electrical activity observed in slime moulds, plants, and fungi, I will demonstrate how simple biological networks exhibit complex decision-making and adaptive behaviours – without neurons or brains. These organisms challenge our traditional notions of cognition, revealing that intelligence may not require a nervous system at all. I will also present recent work from our laboratory where we build and study proto-brains: experimental ensembles of proteinoid microspheres that spontaneously generate spiking patterns, coordinate actions, and process environmental information. By examining these synthetic and natural minimal systems, we gain insights into the deep origins of intelligence, suggesting that cognition may emerge wherever matter organises to sense, decide, and act upon the world.



**Prof. Dr. Selmer
Bringsjord**
Rensselaer

Polytechnic Institute,
Rensselaer Artificial
Intelligence and
Reasoning Laboratory,
New York, USA

Variegated Common Sense Versus the Science of Universal Intelligence, When the Aliens Arrive

The abstract is available at:

<https://kryten.mm.rpi.edu/SBringsjordWhenTheAliensArrive0513250054.pdf>



Dr. David Gamez

Department of
Computer Science at
Middlesex University,
London, UK

Intelligence and Consciousness in Natural and Artificial Systems

Considerable progress has been made with the development of systems that can drive cars, play games, predict protein folding and generate natural language. These systems are described as intelligent and there has been a great deal of talk about the rapid increase in artificial intelligence and its potential dangers. However, our theoretical understanding of intelligence and ability to measure it lag far behind our capacity for building systems that mimic intelligent human behaviour. There is no commonly agreed definition of the intelligence that AI systems are said to possess, nor has anyone developed a practical measure that would enable us to compare the intelligence of humans, animals and AIs on a single scale. This talk addresses these problems by clarifying the nature of intelligence and outlining a new algorithm for measuring intelligence that can be applied to any system.

The first part of the talk starts with a discussion of previous definitions of intelligence. It then argues for a close link between prediction and intelligence and addresses two misconceptions about intelligence. The first is that people often think that humans have a general form of intelligence that has the same level in all environments. This belief motivates the idea that we could develop machines with artificial general intelligence (AGI). However, human intelligence often fails when it is confronted with environments that are significantly different from the natural world, such as high-dimensional numerical spaces. A second issue is that people naively assume that they directly apply their intelligence to the physical world. However, we can only be intelligent about things that are revealed to us through our senses, and people, animals and artificial systems have very different sensory experiences. So, it is much more accurate to say that agents apply their intelligence to their perceived environment, or *umwelt*.

The second part of the talk explores the measurement of intelligence. Previous

work in this area includes IQ, g and universal measures, such as compression tests and algorithms based on goals and rewards. To address the limitations of previous measures, I have developed a new algorithm for measuring predictive intelligence that is based on an agent's internal state transitions. Experiments have been done to test this algorithm, and it has many potential applications in AI safety and the comparative study of intelligence.

The talk concludes with some reflections on the relationships between intelligence and consciousness. It is commonly assumed that there is a close relationship between intelligence and consciousness in biological systems, However, this correlation might not exist in artificial systems, who could be highly intelligent with low levels of consciousness, or highly conscious with low levels of intelligence. In the future we might be able to use algorithmic measures of consciousness, such as information integration theory (IIT), and universal measures of intelligence to systematically study the relationships between intelligence and consciousness in natural and artificial systems.

Keywords: Intelligence and Consciousness in Natural and Artificial Systems



Prof. Dr. Michael Levin

Levin Lab, School of Arts and Sciences.
Department of Biology,
Tufts University,
Massachusetts, USA

Mind everywhere: recognizing and communicating with unconventional intelligence

In this talk I will describe my framework for an engineering approach to recognizing, communicating with, and ethically relating to unconventional intelligences. I will begin with some conceptual clarification of intelligence and embodiment. I will then describe a number of surprising examples of intelligence, using the collective intelligence of cells navigating anatomical space as a model system, and show data on our approach to use the bioelectric interface as a means of communicating with the agential material of life. I will show how re-setting the memories, goals, and cognitive light cone of cell groups drives applications in birth defects, regenerative medicine, cancer, and bioengineering. I will describe our efforts to develop tools to expand our native mind-blindness to large regions of the cognitive spectrum, including the emergent cognition of very minimal models (both living and computational) whose capabilities are not explained by evolutionary history. I will end with some speculative implications of these ideas for the future of science and philosophy of mind.



Prof. Dr. Lorenzo Magnani

Professor, Department of Humanities,
Philosophy Section
University of Pavia,
Pavia, Italy (Ret.)

Abductive Intelligence and Creativity-The Role of Eco-Cognitive Openness and Situatedness

I will use my studies on abductive intelligence in an eco-cognitive framework to demonstrate the concept of "locked and unlocked strategies" in deep learning systems, indicating different inference routines for creative results. Higher forms of creative abductive intelligence are involved in unlocked human cognition, whereas locked abductive strategies are characterized by weak hypothetical creative intelligence because of the absence of what I refer to as eco-cognitive openness and situatedness. The fundamental nature of the human brain as an open system that is continuously coupled with the environment - a so-called "open" or dissipative system - is the physical basis for this special type of "openness". The brain's activity is the continuous attempt to achieve equilibrium with its environment, and this interaction can never be turned off without seriously harming the brain. It is impossible to imagine the brain lacking its physical essence, which is its openness. In the brain, ordering is the direct result of an "internal" open dynamical process of the system rather than being generated from the outside, as I have described in my latest book *Eco-Cognitive Computationalism* (2022), "computational domestication of ignorant entities".

No Intelligence without Representation

This talk focuses on the indispensable role of representation in constituting intelligence. Central to the inquiry into intelligence is understanding the relationship between mind, world, and action. I argue, first, that even basic embodied intelligence, often cited as evidence against representation, relies fundamentally on contentful states. Drawing on analyses of intentionality and procedural memory, skills are shown to possess directive content (a world-to-mind direction of fit) defined by satisfaction conditions. This representational



Assoc. Prof. Marcin

Milkowski

Associate professor in
the Institute of
Philosophy and
Sociology of the Polish
Academy of Sciences,
Warsaw, Poland

layer is crucial for explaining the normativity, learning, and potential failures (e.g., apraxia) inherent in skilled action. Second, I argue that directive content alone is insufficient. Genuine intelligence requires the capacity for descriptive representation—states with a mind-to-world direction of fit, capable of being true or false in a way that corresponds to reality. This commitment to representational realism and truth is not merely a feature of high-level cognition but a foundational requirement for systems that can learn about, understand, and flexibly adapt to the complexities of the world beyond immediate interaction loops. Intelligence, therefore, is inextricably linked to the dual representational capacities to both shape the world according to goals and accurately reflect its state.



Prof. Dr. Vincent C. Müller

A. v. Humboldt
Professor, Philosophy
& Ethics of AI,
University of Erlangen-
Nürnberg, Germany

AI Philosophy: A New Philosophical Method (Vincent C. Müller with Guido Löhr)

On the background of the general problem of philosophical methodology, we identify a new tool for the philosophical toolbox: AI. We propose that not only can AI learn from philosophy, but philosophy can learn from AI, too: It is both ways. This applies particularly to conceptual analysis, which can be advanced by asking what would be required for an AI system to fall under the concept we are discussing. This method it avoids anthropocentrism and gives us a new way of testing our philosophical theories. Given the wide range of features we can consider for AI systems, this method allows us to cover a wide range of philosophical issues, especially in the philosophy of mind, language, epistemology, and ethics. In the first section, we present two salient examples of this new method (consciousness and free will) and in the second section, we analyse what the method is. Finally, we defend this new method against anticipated objections.



Prof. Dr. Diane Proudfoot

Professor of
Philosophy, University
of Canterbury - TE
WHARE WĀNANGA O
WAITAHA, New
Zealand

Four myths about Turing's view of intelligence and his test

In current inquiry into intelligence, Turing's work is regarded as foundational. Yet misunderstandings of his view of (the concept of) intelligence and his famous test of intelligence in machines are widespread. I shall argue that four standard interpretations of Turing and his imitation game are mistaken. First, that Turing was (in an influential sense) a behaviourist about the mind. Second, that his approach to the mind was (again in an influential sense) computationalist. Third, that Turing's philosophy of mind was radically different from that of his contemporary, Wittgenstein, who in turn was a severe critic of Turing. And last, that Turing's test has been passed by recent GPT models—with the result that we need a new test of intelligence in machines.



Prof. Dr. Oron Shagrir

Vice-President for
International Affairs;
Schulman Chair in of
Philosophy, The
Hebrew University of
Jerusalem, Israel

The mathematical objection to artificial (machine) intelligence

Turing develops the idea of machine intelligence in a series of lectures and papers between 1947 and 1952. In some of them he addresses the mathematical objection (his term) whose gist is the claim that humans can assert some mathematical truths that exceed the abilities of computing machines. We first ask why Turing took so seriously the mathematical objection. After all, even if some humans surpass machines in their mathematical abilities, this by itself does not undermine the project of machine intelligence. Our answer is that the mathematical objection raises a dilemma with respect to Turing's core claims about machine intelligence and forces him to relinquish at least one of them. We then clarify Turing's reply to the mathematical objection. Based on the textual evidence, we argue that, according to Turing, the machine that plays against the human in the Turing test is not a static machine but an enhanced machine. Keywords: The mathematical objection to artificial (machine) intelligence



**Prof. Dr. Jordi
Vallverdú**

Professor & ICREA
Acadèmia researcher,
Universitat Autònoma
de Barcelona,
Catalonia, Spain

Cognitive Romanticism: Humans Are Worse than Stochastic Parrots!

Most Humans Are Stochastic Parrots, and LLMs Reveal Our Intellectual Mediocrity. While critics often dismiss large language models (LLMs) as mere "stochastic parrots," I argue that this accusation misunderstands both machine intelligence and, more importantly, human cognition itself. Most human thought is not creative, critical, or unpredictable; it is rote, imitative, and driven by deeply ingrained social, religious, and cognitive biases. The dominance of myth, ideology, and irrational belief systems across cultures reveals that humans themselves function largely as stochastic parrots – endlessly repeating patterns they neither question nor understand. Rather than exposing the limitations of artificial intelligence, LLMs expose the uncomfortable truth about human intellectual mediocrity. In this talk, I will attack the myth of human cognitive exceptionalism, dismantle the romantic notions attached to human "creativity" and "autonomy," and propose a radical redefinition of intelligence beyond humanist illusions. Intelligence, whether in biological or artificial systems, must be seen not as the privilege of a superior species, but as an emergent property of patterned interaction with environments – often stochastic, occasionally innovative, but rarely transcendental.

Keywords: Cognitive Romanticism: Humans Are Worse than Stochastic Parrots!



Dr. Hector Zenil

Associate Professor /
Senior Lecturer,
School of Biomedical
Engineering & Imaging
Sciences, Faculty of
Life Sciences &
Medicine & King's
Institute for Artificial
Intelligence, King's
College London, UK

Why LLMs Can't Escape the Pattern-Matching Prison if They Don't Learn Recursive Compression

In this talk we will introduce and discuss SuperARC, a new proposed open-ended test based on algorithmic probability to critically evaluate claims of Artificial General Intelligence (AGI) and Artificial Superintelligence (ASI), challenging standard metrics grounded in statistical compression and human-centric tests. By leveraging Kolmogorov complexity rather than Shannon entropy, the test measures fundamental aspects of intelligence, such as synthesis, abstraction, and planning or prediction. Comparing state-of-the-art Large Language Models (LLMs) against a hybrid neurosymbolic approach, we identify inherent limitations in LLMs, highlighting their incremental and fragile performance, and their optimisation primarily for human-language imitation rather than genuine model convergence. These results prompt philosophical reconsideration of how intelligence—both artificial and natural—is conceptualised and assessed.

Program Overview

10 Jun 2025	11 Jun 2025	12 Jun 2025	13 Jun 2025	14 Jun 2025
AI for Philosophy	Intelligence Beyond Anthropomorphization: Intelligence and Life	Critical Discussions of AI, Values, Ethics	Intelligence Beyond Anthropomorphization: Widening the Perspective	From Human to Artificial Intelligence

IOCPH 2025 Program

10 June 2025 (Tuesday)

AI for Philosophy

Time: 12:00 (CEST, Basel) / 06:00 (EDT, New York) / 18:00 (CST Asia, Beijing)

Time (CEST)	Type	Speaker	Title
12:00–12:30	Conference Opening	Marcin J. Schroeder & Gordana Dodig Crnkovic Conference Chairs	Welcome speech from the Conference Chairs
12:30–13:00	Contributed paper	Zijian Ding School of Philosophy, Psychology and Language Sciences, University of Edinburgh, UK	Intelligence as Typological Cognition: Revisiting Jungian Functions for Human and Artificial Minds
13:00–14:00	Invited Speaker	Vincent C. Müller	AI Philosophy: A New Philosophical Method (Vincent C. Müller with Guido Löhr)
		Break from 14:00–15:00	

Time (CEST)	Type	Speaker	Title
15:00–15:30	Contributed paper	Marcin J. Schroeder	Intelligence as the Capacity to Overcome Complexity of Information: Search for Unity in the Diverse Forms of Intelligence
15:30–16:00	Contributed paper	Jaime F. Cárdenas-García University of Maryland, Baltimore County, USA	Infoautopoiesis and Intelligence
16:00–17:00	Invited Speaker	Lorenzo Magnani	Abductive Intelligence and Creativity-The Role of Eco-Cognitive Openness and Situatedness
		Break from 17:00–18:00	

Time (CEST)	Type	Speaker	Title
18:00–19:00	Keynote Speaker	Stephen Wolfram	Computational Foundations of Minds and the Universe
19:00–20:00	Keynote Speaker & Moderator	Stephen Wolfram & Gordana Dodig Crnkovic	Keynote discussion
20:00–21:00	Panel Discussion	Edward A. Lee Jaime Cardenas-García Lorenzo Magnani Sheri Markose Vincent C. Müller	Panel 1: AI for Philosophy

11 June 2025 (Wednesday)

Intelligence Beyond Anthropomorphization: Intelligence and Life

Time: 9:00 (CEST, Basel) / 03:00 (EDT, New York) / 15:00 (CST Asia, Beijing)

Time (CEST)	Type	Speaker	Title
09:00–09:30	Contributed Paper	Gordana Dodig Crnkovic	Evolution of Intelligence from Active Matter to Complex Intelligent Systems: An Agent-Based Approach

09:30–10:00	Contributed Paper	Paweł Polak Pontifical University of John Paul II in Krakow, Faculty of Philosophy, Poland	Towards Beneficial AI: Biomimicry framework to design intelligence cooperating with biological entities
10:00–10:30	Contributed Paper	Nina Poth Radboud University Nijmegen, Nijmegen, Netherlands	Intelligent behaviour as adaptive control guided by accurate prediction
10:30–11:00	Contributed Paper	Bartosz Michał Radomski Ruhr University Bochum, Bochum, Germany	Three Puzzles of Adaptivity: A Lens to Understand Definitions of Life, Cognition, and Intelligence
11:00–12:00	Invited Speaker	Andrew Adamatzky	Origin of Intelligence
		BREAK: 12:00–13:00 CEST	

Time (CEST)	Type	Speaker	Title
13:00–14:00	Invited Speaker	Michael Levin	Mind everywhere: recognizing and communicating with unconventional intelligence
14:00–15:00	Panel Discussion	Andrew Adamatzky Edward A. Lee Lorenzo Magnani Michael Levin	Panel 2 Intelligence Beyond Anthropomorphization
15:00–15:30	Contributed paper	Kyoko Nakamura Osaka University, Osaka, Japan	Natural Born Intelligence that summons <emotions=politics>
15:30–16:00	Contributed paper	Kyoko Nakamura Osaka University, Osaka, Japan	Body as Anti-Anthropomorphic Landscape: Natural Born Intelligence in Bog Body
16:00–16:30	Contributed paper	Soumya Banerjee University of Cambridge, UK	Cosmicism and Artificial Intelligence: Beyond Human-Centric AI
16:30–17:00	Contributed paper	Jevgenija Sivoronova Daugavpils University, Daugavpils, Latvia	The Cognition System Theory
17:00–17:30	Contributed paper	Maria Olon Tsaroucha Supraconscious You, NY, USA	Beyond the Fifth Wall: Acting, Avatars, and the Quantum Self—Toward a Unified Framework for Human Intelligence and Identity
17:30–18:00	Posters	1.) Under the reign of AI: a new 21st century taste dispute? Reflections about Ian Hamilton Finlay 'Little Sparta' and John Goto 'High Summer'. - Maria De Fátima Alexandrino Alves de Sá Monteiro Lambert 2.) The dialectical philosophy of intelligent model and mathematical physics simulation in the wave motion research - Peng Ren 3.) When Planes Fly Better Than Birds: Should AIs Think Like Humans? - Soumya Banerjee 4.) Thinking as Decolonial Praxis - Colette Sybille Jung 5.) Jean Piaget and Objectivity—Genetic Epistemology's Place in a View from Nowhere - Mark A Winstanley 6.) Key Methodologies in Intelligence Studies: Techniques, Skills, and Integration - Sunila Atul Patil	FLASH POSTER PRESENTATION

		7.) The concept of Information and the Transmission of the Experience of Beauty - Łukasz Mściśławski	
--	--	--	--

12 June 2025 (Thursday)***Critical Discussion of (Generative) AI******Time: 9:00 (CEST, Basel) / 03:00 (EDT, New York) / 15:00 (CST Asia, Beijing)***

Time (CEST)	Type	Speaker	Title
09:00–09:30	Contributed paper	Marcus Abundis Formerly Stanford University (March 2011), USA	Intelligence and The “Hard Problem” of Consciousness – a short analysis
09:30–10:00	Contributed paper	Saskia Janina Neumann Eötvös Lorand University, Budapest, Hungary	How Brooks' behaviour based robots teach us a lesson about the definition of knowledge
10:00–10:30	Contributed paper	Rafal Maciag Jagiellonian University, Institute of Information Studies, Kraków, Poland	Knowledge as an emergent effect of the complexity of large language models
10:30–11:00	Contributed paper	Laida Arbizu Aguirre PhD Student at the University of the Basque Country (UPV/EHU), Biscay, Spain	Epistemic Architectures of Intelligence: A Feminist Critique of Androcentric Reason
11:00–12:00	Invited Speaker	Oron Shagrir	The mathematical objection to artificial (machine) intelligence
		Break from 12:00–13:00	

Time (CEST)	Type	Speaker	Title
13:00–14:00	Invited Speaker	Hector Zenil	Why LLMs Can't Escape the Pattern-Matching Prison if They Don't Learn Recursive Compression
14:00–15:00	Invited Speaker	Jordi Vallverdú	Cognitive Romanticism: Humans Are Worse than Stochastic Parrots!
15:00–16:00	Panel Discussion	Hector Zenil, Selmer Bringsjord, Jordi Vallverdú, David Gamez	Panel 3 Critical Discussion of (Generative) AI
16:00–16:30	Contributed paper	José Ferraz-Caetano LAQV-REQUIMTE Faculty of Sciences, University of Porto, Portugal	Epistemic Automation: Using Machine Learning to Scale and Interpret Agent-Based Models of Scientific Inquiry
16:30–17:00	Contributed paper	Szymon Miłkoś Polish Academy of Sciences, Poland	Proto-Intelligent Inquiry
17:00–17:30	Contributed paper	Maria Regina Brioschi University of Milan, Italy	From GenAI to Human Beings and Back: Assessing Creative Intelligence
17:30–18:00	Contributed paper	Johan Largo Institute of Philosophy, University of Luxembourg, Luxembourg	Rethinking intelligence: the problems of the representational view in the era of LLMs

13 June 2025 (Friday)***Intelligence Beyond Anthropomorphization: Widening the Perspective***

Time: 9:00 (CEST, Basel) / 03:00 (EDT, New York) / 15:00 (CST Asia, Beijing)

Time (CEST)	Type	Speaker	Title
09:00–09:30	Contributed paper	Jian Wang Xi'an Jiaotong University, Xi'an, China	A Framework for the Ethical Design and Accountability Attribution of Military Robot Systems
09:30–10:00	Contributed paper	FEI SUN Computer Science and Engineering, Mälardalens University, Sweden	Axiology and the Evolution of Ethics in the Age of AI
10:00–10:30	Contributed paper	Attila Egri-Nagy Akita International University, Akita, Japan	Thought Structures and their Morphisms
10:30–11:00	Contributed paper	Anna Sarosiek University of the National Education Commission, Krakow, Poland	A Cybernetic Approach to the Intentionality of Artificial Systems: A Regulatory Function Instead of a Representational One?
11:00–12:00	Invited Speaker	David Gamez	Intelligence and Consciousness in Natural and Artificial Systems
		Break from 12:00–13:00	

Time (CEST)	Type	Speaker	Title
13:00–14:00	Invited Speaker	Marcin Milkowski	No Intelligence without Representation
14:00–15:00	Invited Speaker	Selmer Bringsjord	Variegated Common Sense Versus the Science of Universal Intelligence, When the Aliens Arrive
15:00–16:00	Panel Discussion	Selmer Bringsjord, Oron Shagrir, David Gamez	Panel 4: Intelligence Beyond Anthropomorphization: Widening the Perspective
16:00–16:30	Contributed paper	Velina Slavova New Bulgarian University, Sofia, Bulgaria	Transitive Self-Reflection – A Fundamental Criterion for Detecting Intelligence
16:30–17:00	Contributed paper	Sheri Markose University of Essex, Essex, UK	Gödelian Self-Referential Genomic Information Processing: Complex Self-Other Interactions, Social Cognition and Novelty Production
17:00–17:30	Contributed paper	Jerry LR Chandler George Mason University, Fairfax, USA	Biological Information Theory: Group Identity Nexuses of Formal Grammars, Numbers, and Propositions
17:30–18:00	Contributed paper	Rao Mikkilineni Golden Gate University, San Francisco, USA	Application of Burge's General Theory of Information: Autopoietic and Meta-Cognitive Machines

14 June 2025 (Saturday)

From Human to Artificial Intelligence

Time: 9:00 (CEST, Basel) / 03:00 (EDT, New York) / 15:00 (CST Asia, Beijing)

Time (CEST)	Type	Speaker	Title
09:00–09:30	Contributed paper	Daniel Boyd Independent researcher	Comparison of the effect of language on high level information processes in humans and linguistically mature generative AI
09:30–10:00	Contributed paper	Angelo Complerchio Luleå University of	Free Intelligence in the Wild: on Human and other Non-human Natural and Artificial forms

		Technology, Norrbotten County, Sweden	
10:00–11:00	Invited Speaker	Diane Proudfoot	Four myths about Turing's view of intelligence and his test
11:00–12:00	Panel Discussion	Diane Proudfoot Lorenzo Magnani	Panel 5: From Human to Artificial Intelligence
		Break from 12:00–13:00	

Time (CEST)	Type	Speaker	Title
13:00–13:30	Contributed paper	Liang Wang Xi'an Jiaotong University, China	Ethical Design of Social Robots Based on Confucian Etiquette Practices
13:30–14:00	Contributed paper	Francisco Miguel Macías-Pozo University of Málaga, Spain	Intelligence as a second-order virtue: changing attitudes for successful interactions in digital environments
14:00–14:30		BREAK	
14:30–15:00	Contributed paper	Alexander Bringsjord Rensselaer Polytechnic Institute (RPI), New York, USA	The Threat to Human Intelligence From AI of Today
15:30–16:00	Contributed paper	Dorothea Olkowski University of Colorado, Colorado Springs, CO USA	You Think You Want a Revolution?
16:00–16:30		Marcin J. Schroeder & Gordana Dodig Crnkovic Moderators	Closing plenary discussion

General Session

sciforum-123394: The Threat to Human Intelligence From Today's AI

Alexander Bringsjord

Rensselaer Polytechnic Institute (RPI)

Reflections upon the nature of today's Artificial Intelligence (AI), on the one hand, and human intelligence (HI), on the other, are now intertwined, catalyzed by the stunning increase in the breadth and speed of artificial agents produced (mostly) by for-profit corporations, and used by the many humans who pay for the privilege of using (the larger, more intelligent(?) versions of) these agents. Unfortunately, because of the particular nature of this AI, this reflection constitutes an attack on HI. This is so because the basis for today's ascension of AI is a form of venerated machine learning (ML) inconsistent with a large part of what, undeniably, distinguishes HI, and has done so since the dawn of recorded history. Specifically, the sub-type of ML known as "deep learning" (DL) is fast reducing the received conception of HI to the sub-human level, while at the same time audaciously raising the banner of "foundation models" over today's AI to signal the subsuming of all of HI (see, e.g., Agüera Y Arcas, B. & Norvig 2023). Humans have long been distinguished by their ability to create and assess chains of logical reasoning—deductive, analogical, inductive, abductive—over internally stored, structured, declarative knowledge, in order to produce conclusions, themselves structured and declarative. But AI qua DL literally has no such knowledge in the first place: there is nowhere in an artificial deep neural network (DNN) where such knowledge resides (see, e.g., Russell & Norvig 2020; and for a lively, lucid treatment of the issue, see Wolfram 2023). Ironically, none of the corporations such as OpenAI that sell and promote DL could survive for a week without humans therein employing logical reasoning on a broad scale. Genuine HI, for a simple but telling example, has for millennia consisted in part in being able to (i) memorize the algorithm of long multiplication over two whole numbers, (ii) apply it to compute the function of multiplication over, for instance, 23 and 5 (to yield 115), and (iii) certify as correct and gauge, logically, how feasible this algorithm is in the general case. AI qua DL does not have even this intelligence, since no algorithms whatsoever are stored in a DNN; and as a matter of settled science, DNNs only approximate the computing of functions. Hence, since HI is increasingly identified with DL (and human minds identified with DNNs), HI is regarded to be dramatically lower than what it is. In short, today's DNN-based GenAI, itself painfully illogical (see, e.g., Arkoudas 2023a, 2023b), portrays HI as constitutionally illogical. This is a growing cancer that philosophy as a field must uncover, and seek to heal, or at least mitigate. Our ability to engage in logical reasoning is in large part what sets us apart from all other natural species, and this hallmark, while under attack today, must be protected. The need to do so is all the more dire because Artificial General Intelligence (AGI), taken to have HI as its initial target (before aiming at superhuman intelligence), is defined as being devoid of logical reasoning (see, e.g., the entirely statistical Legg & Hutter 2007).

I end by (briefly) offering and defending a course of action to address the threat to HI.

References

- Agüera Y Arcas, B. & Norvig, P. (2023) "Artificial General Intelligence is Already Here" Noema, October.
- Arkoudas, K. (2023a) "ChatGPT is No Stochastic Parrot. But It Also Claims That 1 is Greater than 1" Philosophy & Technology 36.3: 1–19.
- Arkoudas, K. (2023b) "GPT-4 Can't Reason: Addendum" Medium.
https://medium.com/@konstantine_45825/gpt-4-cant-reason-addendum-ed79d8452d44
- Russell, S. & Norvig, P. (2020) Artificial Intelligence: A Modern Approach (New York, NY: Pearson).
- Legg, S. & Hutter, M. (2007) "Universal Intelligence: A Definition of Machine Intelligence" Minds and Machines 17.4: 391–444.

Wolfram, S. (2023) "What is ChatGPT Doing ... and Why Does It Work?"

<https://writings.stephenwolfram.com/2023/02/what-is-chatgpt-doing-and-why-does-it-work>

Philosophical Issues

Such issues are myriad in the present case. The nature of intelligence—artificial, natural/human, supernatural/superhuman—must be engaged. A brief intellectual history, rooted in philosophy, must be provided for context. The growing “science of intelligence” within the fields of AI and AGI must be analyzed philosophically. The analysis of how humans learn must be compared with the analysis of how ML (DL in particular; and RL) systems “learn.” Substantiation as to how today’s “GenAI” seriously limits intelligence in both AI and HI, and threatens the latter, must be provided by argument. The presentation/paper concludes with a proposed solution for improving the current state of affairs.



© 2025 by the author(s). Licensee MDPI, Basel, Switzerland. This article is an open access article distributed under the terms and conditions of the Creative Commons Attribution (CC BY) license (<http://creativecommons.org/licenses/by/4.0/>).

sciforum-124987: Four myths about Turing's view of intelligence and his test

Diane Proudfoot

University of Canterbury, Christchurch, New Zealand

In current inquiry into intelligence, Turing's work is regarded as foundational. Yet misunderstandings of his view of (the concept of) intelligence and his famous test of intelligence in machines are widespread. I shall argue that four standard interpretations of Turing and his imitation game are mistaken. First, that Turing was (in an influential sense) a behaviourist about the mind. Second, that his approach to the mind was (again in an influential sense) computationalist. Third, that Turing's philosophy of mind was radically different from that of his contemporary, Wittgenstein, who in turn was a severe critic of Turing. And last, that Turing's test has been passed by recent GPT models—with the result that we need a new test of intelligence in machines.



© 2025 by the author(s). Licensee MDPI, Basel, Switzerland. This article is an open access article distributed under the terms and conditions of the Creative Commons Attribution (CC BY) license (<http://creativecommons.org/licenses/by/4.0/>).

sciforum-124091: Mind everywhere: recognizing and communicating with unconventional intelligence

Michael Levin

School of Arts and Sciences. Department of Biology, Tufts University, Massachusetts, USA

In this talk I will describe my framework for an engineering approach to recognizing, communicating with, and ethically relating to unconventional intelligences. I will begin with some conceptual clarification of intelligence and embodiment. I will then describe a number of surprising examples of intelligence, using the collective intelligence of cells navigating anatomical space as a model system, and show data on our approach to use the bioelectric interface as a means of communicating with the agential material of life. I will show how re-setting the memories, goals, and cognitive light cone of cell groups drives applications in birth defects, regenerative medicine, cancer, and bioengineering. I will describe our efforts to develop tools to expand our native mind-blindness to large regions of the cognitive spectrum, including the emergent cognition of very minimal models (both living and computational) whose capabilities are not explained by evolutionary history. I will end with some speculative implications of these ideas for the future of science and philosophy of mind.



© 2025 by the author(s). Licensee MDPI, Basel, Switzerland. This article is an open access article distributed under the terms and conditions of the Creative Commons Attribution (CC BY) license (<http://creativecommons.org/licenses/by/4.0/>).

sciforum-119545: Axiology and the Evolution of Ethics in the Age of AI

Fei Sun

Computer Science and Engineering, Mälardalens University

Artificial intelligence (AI), particularly autonomous systems, challenges traditional ethical frameworks by reshaping human values, agency, and responsibility. This paper argues that axiology—the philosophical study of values—offers a critical foundation for AI ethics by accommodating the dynamic relationship between technology and morality. Unlike rigid ethical theories, axiology provides an adaptive approach to algorithmic bias, depersonalized healthcare, and AI-mediated governance.

We propose an integrative axiological model that synthesizes deontological, utilitarian, and virtue ethics to ensure that AI aligns with pluralistic human values. This framework balances duty (transparency, fairness), outcomes (social good, efficiency), and virtue (human dignity, trust), akin to multicriteria decision analysis (MCDA) (Sapienza, Dodig-Crnkovic, Crnkovic, 2016), which systematically evaluates competing priorities in complex decision-making. For example, while utilitarianism might favor AI's cost-saving healthcare diagnostics, virtue ethics ensures patient autonomy remains central, and deontology requires transparency in algorithmic decisions. This synthesis prevents AI from privileging one value (e.g., efficiency) at the expense of others (e.g., privacy) but relies on multicriterion-based decisions.

Case studies illustrate AI's dual impact: In education, AI-powered learning enhances accessibility but may risk dehumanizing assessment. In healthcare, AI-driven diagnostics improve accuracy, yet excessive reliance on AI may threaten patient trust if empathy is overlooked. In governance, AI improves transparency but may raise ethical concerns over surveillance and bias in policing. These examples underscore the need for an evolutionary ethics, where values shift alongside technological advances.

This model aligns with Digital Humanism, which resists reducing humans to data points, and Responsible AI, which prioritizes accountability. Together, they advocate for AI that enhances—not undermines—human dignity, equity, and democratic agency.

To prevent ethical stagnation, policymakers and developers may adopt this axiological lens, ensuring that AI evolves as a tool for societal flourishing rather than a destabilizing and depersonalized force. By focus on axiology, we reframe AI ethics as a living discipline—one that reconciles competing values and safeguards humanity's moral commitments in an AI algorithmic age.

References

1. Sapienza, G., Dodig-Crnkovic, G. and Crnkovic, I. (2016) "Inclusion of Ethical Aspects in Multi-Criteria Decision Analysis". Proc. WICSA and CompArch conference. Decision Making in Software ARCHitecture (MARCH), 2016 1st International Workshop. Venice April 5-8, 2016. DOI: 10.1109/MARCH.2016.5, ISBN: 978-1-5090-2573-2. IEEE



© 2025 by the author(s). Licensee MDPI, Basel, Switzerland. This article is an open access article distributed under the terms and conditions of the Creative Commons Attribution (CC BY) license (<http://creativecommons.org/licenses/by/4.0/>).

sciforum-125005: Computational Foundations of Minds and the Universe

Stephen Wolfram^{1,2,3,4}

¹ Founder & CEO of Wolfram Research

² Creator of Mathematica, Wolfram|Alpha & Wolfram Language

³ Author of A New Kind of Science and other books

⁴ Originator of Wolfram Physics Project

Wolfram's recent work on the foundations of physics, mathematics, biology and machine learning introduces a major new framework for thinking about fundamental philosophical questions. This talk will provide a non-technical, philosophically oriented survey of these directions.

<https://writings.stephenwolfram.com/category/philosophy/>



© 2025 by the author(s). Licensee MDPI, Basel, Switzerland. This article is an open access article distributed under the terms and conditions of the Creative Commons Attribution (CC BY) license (<http://creativecommons.org/licenses/by/4.0/>).

sciforum-119923: Intelligence as the Capacity to Overcome the Complexity of Information: The Search for Unity in Diverse Forms of Intelligence

Marcin J. Schroeder

Faculty of International Liberal Arts, Akita International University, 193-2 Okutsubakidai, Tsubakigawa, Yuwa, Akita-shi, 010-1211 Akita, Japan

1. Introduction and Motivation

In his 2011 bestseller book “*Thinking Fast and Slow*” [1], Daniel Kahneman introduced one of his rules of common thinking: “A reliable way to make people believe in falsehoods is frequent repetitions because familiarity is not easily distinguished from truth.” This rule applies to sentences that can be qualified as true or false. A similar rule can be applied to concepts linking words or terms with their meanings. There seems to be sociolinguistic correspondence between the frequency of the use of terms or expressions, in particular those with philosophical significance, and the unrecognized diversity of their understanding. This diversity is obscured by the fact that the frequency of the use of these words generates the impression of familiarity and familiarity the illusion of an identity and existence of a unique denotation independent from any inquiry, the conviction that the meaning is obvious and uniform. As a result, people use frequently invoked words with a diverse subjective understanding while being convinced about their objective, uniform, commonly shared meaning. If they are confronted with differences in the understanding presented by others, they claim that such a different understanding is erroneous.

Intelligence, either artificial, natural, or human, has become an illustrative instance of this regularity. The term is present everywhere, in particular when qualified as artificial and used in its staple abbreviation AI, or when it appears unqualified when understood as human.

There are some curious differences between the ways in which human and artificial forms of intelligence are viewed. For a long time, there has been an unusual consensus, among experts and laypeople, that human intelligence is not only diverse but that its different forms are independent and possibly uncorrelated, each with a separate psycho-neurological mechanism. Thus, human intelligence can be fluid or crystallized following the division introduced in 1943 by Raymond Cattell [2], splitting Charles Spearman’s concept of general intelligence present in psychology since the beginning of the 20th century into two different capacities. The former was understood as a purely general ability to solve unexpected problems without prior preparation or experience, and the latter consisted of long-established discriminatory habits acquired through learning or training. Incidentally, this distinction can be considered a precedent to Kahneman’s popular and recent distinction of fast (habitual, rigid) and slow (goal-oriented, flexible) thinking [1].

Further divisions of intelligence came in the 1980s. Robert Sternberg introduced his triarchic theory of intelligence, separating it into analytical, creative, and practical intelligence [3]. At about the same time, Howard Gardner introduced in his 1983 book “*Frames of Mind: Theory of Multiple Intelligences*” [4] the idea of “intelligences” in the plural form and presented a list originally consisting of seven types: linguistic, logical-mathematical, spatial, musical, bodily-kinetic, interpersonal, and intrapersonal. In 1995, he added an eighth type: naturalistic intelligence. This triggered multiple attempts to distinguish specific types of such multiple intelligences. Divisions of human intelligence have been made using diverse criteria and different levels of argumentation, but they have always been motivated by criticism of the inadequacy of the Spearman’s original concept of general intelligence and attempts to measure it using variations of William Stern’s and Alfred Binet’s/Theodore Simon’s IQ tests. Today, there is no agreement regarding the selection of criteria, names, and the degree of correlation, but the idea of multiple intelligences became a standard in the study of human cognitive abilities.

The typical view of artificial intelligence (AI) relates it to human intelligence (such as a simulation, emulation, or recreation of the latter), and the main goal of current technological

research and innovation is the achievement of artificial general intelligence (AGI), which does not have a commonly accepted or even discussed definition but is presented to the general audience in descriptions of the following type: “AGI [...] a system that is capable of matching or exceeding human performance across the full range of cognitive tasks.” [5]. This category error, blurring the distinction between the abstract concept and its individual realizations, characteristic of virtually all instances of the discourse on artificial intelligence (general or otherwise), perpetuates the hidden comprehension of its unity.

There are separate names for specific types of technological realizations of AI (generative AI based on Large Language Models (LLMs) in neural networks with a deep learning architecture, followed by Large Reasoning Models (LRMs) heavily dependent on external prompts, i.e., minimizing their autonomy; agentic AI, re-engaging some forms of algorithmic computation, still dependent on initial prompts but with increasing autonomy in its operation through auto-reprompting; neuro-symbolic AI with a hybrid architecture; causal AI; etc.) However, all of these qualifications reflect not the conceptual diversity of artificial intelligences but the technological differences in the search for the realization of the same goal of artificial general intelligence (AGI).

The inconsistency between conceptualizations of human multiple intelligences and the uniform idea of artificial general intelligence matching or even surpassing human cognitive abilities is only one of the many manifestations of conceptual chaos in the study of intelligence. The situation becomes even more complex when we consider the extensions of intelligence understood as a characteristic of natural but non-human entities present in diverse forms of life on Earth or the expected but not yet known forms of extraterrestrial life. This complication arises not only from the association of intelligence with life, whose conceptualization has been similarly convoluted, but also from the increased diversity of the morphological and behavioral forms involved in manifestations of intelligence in living objects, which require the careful avoidance of anthropomorphization.

2. The Philosophical Significance of Intelligence in Terms of Information and Complexity

Is this conceptual complexity of intelligence or intelligences a good reason to question the feasibility of, or the justification, for any attempts to seek a unifying perspective on such a complex variety of related yet diverse forms of what in multiple contexts is called by the same common-sense name of “intelligence”? This paper has as its objective the justification of a negative answer to this question. In this case, as in many other cases known from the intellectual history of humanity, the phenomenal complexity of the manifestations of intelligence is unquestionable, but the complexity of their perception and comprehension can be overcome with the use of the appropriate intellectual tools.

The tools proposed here are appropriately general concepts of information and its complexity, together with already known methods of reducing or controlling the complexity of information, with their long history going back at least to the Law of Requisite Variety proposed by W. Ross Ashby [6]. In this paper, the methods of reducing complexity are traced much further back in their association with intelligent or efficient inquiry. The use of the concept of information brings its own challenges, as the term information is still a subject of the rule of illusory uniformity of meaning in its common-sense use considered at the beginning of this paper. However, in the philosophy of information, this issue has been adequately addressed, and this intellectual experience can be now applied to intelligence.

Of course, the proposed tools can be used only if we assume that intelligence can be relativized to the concept of information. However, all existing studies of intelligence contain this assumption, even if sometimes in a hidden form. Moreover, the philosophical significance of the concept of information, in particular in its association with complexity, brings into the inquiry an extensive toolbox of philosophical inquiry that allows for the integration of methodologies developed for the research of particular domains of reality. At the same time, the study of intelligence acquires an interdisciplinary status.

After reviewing a wide range of fundamental concepts associated with intelligence in its diverse contexts, this paper presents an initial formulation of the general characterization of intelligence as the ability to minimize the resources necessary to effectively perform a maximal range of actions. However, it is shown that the use of the terms “resource” and “action” leads to overgeneralization in the absence of their qualification or identification of their meaning. For

instance, the motion of every object can be described by the mechanical law of minimal action. On the other hand, the identification of resources may lead either to excessive restrictions, resulting in undergeneralization, or a vicious circle of reasoning. For this reason, the term “resource” is replaced with “information” and “action” with “overcoming complexity”. Then, intelligence is defined as the capacity to overcome complexity. Incidentally, with this understanding of intelligence comes the elimination of the complexity of its manifestations in diverse contexts by lifting the level of abstraction through the overarching concept of information. Therefore, our inquiry can be considered an intelligent inquiry into intelligence.

References

1. Kahneman, D. *Thinking, Fast and Slow*. New York: New York: Farrar, Straus and Giroux, 2011.
2. Cattell, R. B. The measurement of adult intelligence. *Psychological Bulletin*, **1943**, 40(3), 153–193. <https://doi.org/10.1037/h0059973>
3. Sternberg, R. J. *Beyond IQ: A Triarchic Theory of Intelligence*. Cambridge University Press, Cambridge, 1985.
4. Gardner, H. *Frames of Mind: The Theory of Multiple Intelligences*. Basic Books, New York, N.Y. 1983, ISBN 978-0133306149.
5. Jones, N. How AI can achieve human-level intelligence: researchers call for change in tack. *Nature*, News March 4, **2025**, Retrieved March 7, 2025, from: How AI can achieve human-level intelligence: researchers call for change in tack.
6. Ashby, W. R. *An Introduction to Cybernetics*. Chapman & Hall, London, 1956.



© 2025 by the author(s). Licensee MDPI, Basel, Switzerland. This article is an open access article distributed under the terms and conditions of the Creative Commons Attribution (CC BY) license (<http://creativecommons.org/licenses/by/4.0/>).

sciforum-120360: Natural Born Intelligence that summons emotions=politics>

Yukio Pegio Gunji^{1,*} and Kyoko Nakamura²

¹ Intermedia Art and Science, School of Fundamental Science and Engineering, Waseda University

² Osaka University Nakanoshima Art Center, Osaka University

Intelligence and life have been understood in most cases within the framework of artificial intelligence (AI). The framework of AI here is a framework that considers the relationship between two qualitatively different concepts as understanding. Autopoiesis, a model of life, relates the inside and outside of a system in a self-referential way, and the brain, a model of intelligence, is understood as the center that relates various parts to the whole. AI is a type of machine learning that relates unorganized data to organized data. What about swarm intelligence? There too, collective behavior relates individuals to society. All of them are merely restatements of self-reference that relate parts to the whole, and inside to the outside. There is no creativity to go outside these original two concepts.

Emotions can be thought of as politics that suppress individual processing at the lower level through top-down commands. This emergence of emotion = politics does not occur in the framework of AI. Bateson described the gentle bite in animal society as a negation of aggression. An attack is an intense relationship between two logical concepts, enemy and ally. If this were all, it would be nothing more than a variation of self-reference. However, here, there is both an attack and a negation of the attack. If an attack is not presupposed, a soft bite never means a negation of the attack. In other words, a gentle bite is an attack that accepts and relates to both the enemy and the ally and at the same time makes both meaningless. It results in politics=emotion that avoids violence. It was created outside the binary opposition of enemy and ally.

The Natural Born Intelligence (NBI) we are proposing can be considered a generalization of this system of gentle biting, different from a double bind. NBI is an intelligence that simultaneously establishes a positive antinomy that accepts both parties that constitute a binary opposition and a negative antinomy, and summons the outside of the binary opposition. As mentioned above, it can be understood that a simple attack or an AI framework remains a positive antinomy and lacks a negative antinomy.

NBI can be found not only in societies but in brains, animal groups, and even simpler chemical reaction systems. It can be exemplified as a system that uses quantum logic to operate with fluctuations external to quantum logic. The positive antinomy of different Hilbert spaces is quantum entanglement, and ignoring this and accepting fluctuations that go back and forth between different Hilbert spaces is negative antinomy. When both are accepted, quantum logic breaks down, but if one still tries to use quantum logic while maintaining its logicity, then conversely, extremely large fluctuations will be summoned from the outside. They can be thought of as a kind of quantum coherence that resonates with the movements of different Hilbert spaces. When applying the framework of NBI to neurons, there is no doubt that the emergence of emotions = politics can also be understood by using fluctuations external to quantum theory.



© 2025 by the author(s). Licensee MDPI, Basel, Switzerland. This article is an open access article distributed under the terms and conditions of the Creative Commons Attribution (CC BY) license (<http://creativecommons.org/licenses/by/4.0/>).

sciforum-124428: Origin of Intelligence

Andrew Adamatzky

Unconventional Computing Lab, UWE Bristol, UK

What is intelligence, and how deeply is it rooted in the fabric of life itself? In this talk, I explore the fundamental emergence of intelligent behaviours far beyond the animal kingdom. Beginning with spiking electrical activity observed in slime moulds, plants, and fungi, I will demonstrate how simple biological networks exhibit complex decision-making and adaptive behaviours – without neurons or brains. These organisms challenge our traditional notions of cognition, revealing that intelligence may not require a nervous system at all. I will also present recent work from our laboratory where we build and study proto-brains: experimental ensembles of proteinoid microspheres that spontaneously generate spiking patterns, coordinate actions, and process environmental information. By examining these synthetic and natural minimal systems, we gain insights into the deep origins of intelligence, suggesting that cognition may emerge wherever matter organises to sense, decide, and act upon the world.



© 2025 by the author(s). Licensee MDPI, Basel, Switzerland. This article is an open access article distributed under the terms and conditions of the Creative Commons Attribution (CC BY) license (<http://creativecommons.org/licenses/by/4.0/>).

sciforum-123389: Proto-Intelligent Inquiry

Szymon Miłkoś

Polish Academy of Sciences

At the frontier of knowledge, we feel stupid about a given subject, but on the edge of inquiry we also feel stupid about how to proceed—how to choose the next step towards discovery (Schwartz, 2008). This talk introduces *Meta-Creative Problem-Solving* (MCPS), a new concept designed to study and improve inquiry practices and norms as they become intelligent.

To uncover knowledge, we search for *invariances* between epistemic activities (e.g., proposed definitions, models, theories; operationalization of methods; established goals and cultivated values) and the subject of inquiry. The creative part of MCPS operates through directed integration: generating variations within epistemic activities and evaluating consequent changes in relation to the subject of inquiry. The metacognitive dimension of MCPS recognizes that both monitoring and controlling epistemic activities are themselves (meta)epistemic activities. The apparent infinite regress of meta-approaches (i.e., if every epistemic change is guided by another epistemic act, where does it end?) is resolved, i.a., through metacognitive feelings like confidence and fluency (Koriat et al., 2008). These feelings integrate causal information from both present and past problem-solving relationships between epistemic activities and subject of inquiry (Shea, 2023, 2024), guiding complex problem-solving (Ackerman, 2019; Rudolph et al., 2017). They function as half-baked conjectures, preliminary knowledge structures guiding inquiry before formal hypotheses emerge, and on the meta-level—as half-baked heuristics, preliminary methodological structures guiding inquiry before norms or virtues develop.

We don't merely solve problems; we discover how to *do* inquiry itself (Chang, 2022), originating and developing methods, norms, and frames of thinking. The MCPS concept extends Nersessian's (2008) insight that scientific models serve as cognitive artifacts embodying assumptions whose consequences guide model refinement. Similarly, it builds upon Boden's (1991) observation that breakthrough problem-solving often requires reformulating the problem itself. The Meta-Creative Problem-Solving concept generalizes these approaches by recognizing that (meta)epistemic practices and norms themselves embody assumptions whose consequences guide their own refinement. This self-consistency means that during reflection, we may change even the epistemic practices and norms of subsequent reflection, which we can evaluate in light of the subject of inquiry, amid metacognitive feelings crucial at the fractal beginnings of knowledge and inquiry.

The concept of MCPS can help understand—and guide—inquiry practices and norms when no ready-made roadmap exists. This talk will further clarify MCPS through the example of the first mathematical discovery with LLMs (Romera-Paredes et al., 2024) and show how MCPS can guide the development of inquiry norms and practices to integrate theory-driven scientists with data-driven causal discovery algorithms (Andersen, 2024; Petersen et al., 2023; Runge et al., 2023). The concept of Meta-Creative Problem-Solving takes up the challenge of answering how inquiry becomes intelligent.

References:

- Ackerman, R. (2019). Heuristic Cues for Meta-Reasoning Judgments: Review and Methodology. *Psihologijske Teme*, 28(1), 1–20.
- Andersen, H. K. (2024). Why adoption of causal modeling methods requires some metaphysics. *The Routledge Handbook of Causality and Causal Methods*, 87–98.
- Boden, M. A. (1991). *The creative mind: Myths & mechanisms*, Basic Books.
- Chang, H. (2022). *Realism for Realistic People*. Cambridge University Press.
- Koriat, A., Nussinson, R., Bless, H., & Shaked, N. (2008). Information-based and experience-based metacognitive judgments: Evidence from subjective confidence. *Handbook of Metamemory and Memory*, 117–135.

- Nersessian, N. J. (2008). *Creating scientific concepts*. MIT Press.
- Petersen, A. H., Ekstrøm, C. T., Spirtes, P., & Osler, M. (2023). Constructing causal life-course models: Comparative study of data-driven and theory-driven approaches. *American Journal of Epidemiology*, 192(11), 1917–1927.
- Romera-Paredes, B., Barekatin, M., Novikov, A., Balog, M., Kumar, M. P., Dupont, E., Ruiz, F. J. R., Ellenberg, J. S., Wang, P., Fawzi, O., Kohli, P., & Fawzi, A. (2024). Mathematical discoveries from program search with large language models. *Nature*, 625(7995), Article 7995.
- Rudolph, J., Niepel, C., Greiff, S., Goldhammer, F., & Kröner, S. (2017). Metacognitive confidence judgments and their link to complex problem solving. *Intelligence*, 63, 1–8.
- Runge, J., Gerhardus, A., Varando, G., Eyring, V., & Camps-Valls, G. (2023). Causal inference for time series. *Nature Reviews Earth & Environment*, 4(7), 487–505.
- Schwartz, M. A. (2008). The importance of stupidity in scientific research. *Journal of Cell Science*, 121(11), 1771–1771.
- Shea, N. (2024). Metacognition of Inferential Transitions. *Journal of Philosophy*, 121(11), 597–627.
- Shea, N. (2024). *Concepts at the Interface*. Oxford University Press.



© 2025 by the author(s). Licensee MDPI, Basel, Switzerland. This article is an open access article distributed under the terms and conditions of the Creative Commons Attribution (CC BY) license (<http://creativecommons.org/licenses/by/4.0/>).

sciforum-120183: Thought Structures and their Morphisms

Attila Egri-Nagy^{1*} and Miklós Hoffmann²

¹ Akita International University, Japan; egri-nagy@aiu.ac.jp

² Eszterházy Károly Catholic University, Eger, Hungary; hoffmann.miklos@uni-eszterhazy.h

Structure-preserving maps (morphisms) are ubiquitous in mathematics. They are used as analogies and as tools for transforming problems into easier-to-solve formats. We argue that the usefulness of morphisms has a bigger scope. Mathematics has unintentionally developed a grand theory of understanding and explanation since *finding compatible operations across domains is a universal mechanism of any type of intelligence*. Formally, morphisms are described by the equation

$$\varphi(a) * \varphi(b) = \varphi(a \cdot b),$$

where $\cdot$ and $*$ are the two operations in the different domains. We can choose to transfer the inputs and carry out the computation by $*$ or compute with $\cdot$ first, and then transfer the result. If there is a morphic relation, then it does not matter where we perform the operations, as they are compatible.

An archetypical example is a cartographic map, but its geometric or topological features can be generalized. We define *thought structures* using directed graphs where an edge represents passing from one thought to another. The type of connection does not matter much for this all-encompassing definition. It can be a random or guided association, a memory recall, or, in special cases, logical inference. Understanding is then some (partial) morphism between thought structures or other systems. Thinking is compositional (as in category theory), and understanding is morphic (as in algebra). Compatible operations guarantee some useful similarity between the domains.

This argument is not entirely new. Any modeling relationship implies a morphic relation. Conceptual metaphors in cognitive linguistics are also based on the same mechanism. However, there are several misconceptions about the usefulness of morphisms: the seeming rigidity of the mathematical definition (the domains have to be fully defined) ruling out creativity, focusing on isomorphisms (1:1 scale maps) and not taking advantage of information reduction, and missing compositionality by using n -ary relations.

The morphic theory of understanding applies both to *enhancing natural intelligence* by deliberately creating and maintaining morphic relations for thinking, and to *benchmarking artificial intelligence* by looking for morphic representations. In this talk, we will demonstrate how the explicit formulation of morphisms improve explanations and describe these future applications. Further information and references can be found in the preprint at <https://arxiv.org/abs/2411.06806>.



© 2025 by the author(s). Licensee MDPI, Basel, Switzerland. This article is an open access article distributed under the terms and conditions of the Creative Commons Attribution (CC BY) license (<http://creativecommons.org/licenses/by/4.0/>).

sciforum-121130: Why LLMs Can't Escape the Pattern-Matching Prison if They Don't Learn Recursive Compression

Hector Zenil

School of Biomedical Engineering and Imaging Sciences and King's Institute for Artificial Intelligence, King's College London

In this talk we will introduce and discuss SuperARC, a new proposed open-ended test based on algorithmic probability to critically evaluate claims of Artificial General Intelligence (AGI) and Artificial Superintelligence (ASI), challenging standard metrics grounded in statistical compression and human-centric tests. By leveraging Kolmogorov complexity rather than Shannon entropy, the test measures fundamental aspects of intelligence, such as synthesis, abstraction, and planning or prediction. Comparing state-of-the-art Large Language Models (LLMs) against a hybrid neurosymbolic approach, we identify inherent limitations in LLMs, highlighting their incremental and fragile performance, and their optimisation primarily for human-language imitation rather than genuine model convergence. These results prompt philosophical reconsideration of how intelligence—both artificial and natural—is conceptualised and assessed.



© 2025 by the author(s). Licensee MDPI, Basel, Switzerland. This article is an open access article distributed under the terms and conditions of the Creative Commons Attribution (CC BY) license (<http://creativecommons.org/licenses/by/4.0/>).

sciforum-120408: A Framework for the Ethical Design and Accountability Attribution of Military Robot Systems

Jian Wang * and Cong shuai Zhao

Xi'an Jiaotong University

Recent advances in the use of military robots in combat raise serious ethical questions. Military robots have expanded their role in combat, from simple reconnaissance and surveillance to deadly strikes on enemy positions. With the advancement of military robotics technology, the military continues to push for greater autonomy of military robots to reduce the cost of operation and maintenance of military robots. However, as military robots begin to make decisions on their own, the moral responsibilities of military robots become blurred.

If a drone mistakenly destroys a school instead of the right target, who is responsible? Therefore, how to design an ethical military robot? How to reasonably hold a military robot accountable for manslaughter? How the state, society, and even individuals respond to the challenge of military robots will become particularly important and urgent.

This article is divided into five parts: Section 1 provides a background introduction to the ethical design and accountability of military robots. Section 2 summarizes the current status and position of military robots and future development trends, the advantages and disadvantages of military robots compared with humans, the uniqueness of military robot ethics, the difficulty of military robot ethical design, and the dilemma of the accountability of military robots. Section 3 focuses on describing how to design ethical algorithms for military robots. Robotics Ethics outlines and justifies an approach for crafting and assessing ethical algorithms utilized in autonomous machines, including self-driving vehicles and military rescue robots. Derek Leben contends that the evaluation of these algorithms should hinge on their effectiveness in fostering cooperation among self-interested entities. Based on the fact that moral judgments are the product of evolutionary pressures for cooperative behavior in self-interested organisms, this paper compares the results of the application of Rawls' contractualism with the application of utilitarianism, free-willism, deontology, and other moral theories to various dilemma games, and argues that only contractualism can produce Pareto-optimality and cooperative behavior in a variety of dilemmas. The Leximin program, a refined version of the Maximin algorithm based on the contractualist difference principle, is suitable for application to a wide variety of decision spaces. Section 4 centers on the attribution of responsibility for military robots. This paper contends that the challenge of determining the accountability of autonomous robots can be resolved by situating it within the framework of the military chain of command. Decision-making in war is multi-layered, and the military hierarchy is a system that assigns responsibility and limits autonomy among different levels of decision-makers. Section 5 proposes specific countermeasures to address the ethical issues of military robotics. At the technological level, algorithms that continue to optimize the sensitivity of military robots to harm while continuing to increase the level of automation of military robots are an absolute measure to ensure the technological safety of military robots. At the policy and regulatory level, the transparency of military robotics research and development should be strengthened, along with the accountability of those involved.



© 2025 by the author(s). Licensee MDPI, Basel, Switzerland. This article is an open access article distributed under the terms and conditions of the Creative Commons Attribution (CC BY) license (<http://creativecommons.org/licenses/by/4.0/>).

sciforum-124486: Abductive Intelligence and Creativity - The Role of Eco-Cognitive Openness and Situatedness

Lorenzo Magnani

Department of Humanities, Philosophy Section and Computational Philosophy Laboratory, University of Pavia, Pavia, Italy

I will use my studies on abductive intelligence in an eco-cognitive framework to demonstrate the concept of “locked and unlocked strategies” in deep learning systems, indicating different inference routines for creative results. Higher forms of creative abductive intelligence are involved in unlocked human cognition, whereas locked abductive strategies are characterized by weak hypothetical creative intelligence because of the absence of what I refer to as eco-cognitive openness and situatedness. The fundamental nature of the human brain as an open system that is continuously coupled with the environment - a so-called “open” or dissipative system - is the physical basis for this special type of “openness”. The brain's activity is the continuous attempt to achieve equilibrium with its environment, and this interaction can never be turned off without seriously harming the brain. It is impossible to imagine the brain lacking its physical essence, which is its openness. In the brain, ordering is the direct result of an “internal” open dynamical process of the system rather than being generated from the outside, as I have described in my latest book *Eco-Cognitive Computationalism* (2022), “computational domestication of ignorant entities”.



© 2025 by the author(s). Licensee MDPI, Basel, Switzerland. This article is an open access article distributed under the terms and conditions of the Creative Commons Attribution (CC BY) license (<http://creativecommons.org/licenses/by/4.0/>).

sciforum-119567: Application of Burgin's General Theory of Information: Autopoietic and Meta-Cognitive Machines

Rao Mikkilineni

Golden Gate University

Biological systems inherit the knowledge to build, operate, and manage a society of cells with complex organizational structures, where autonomous components execute specific tasks and collaborate in groups to fulfill systemic goals using shared knowledge. This knowledge enables them to receive information through various senses and create a knowledge representation in the form of associative memory and an event-driven interaction history, consisting of various entities, relationships, and behaviors. The General Theory of Information, posited by the late Prof. Mark Burgin, allows us to model knowledge in the form of named sets/Fundamental Triads that capture various entities, relationships, interactions, and behaviors.

In this paper, we describe a digital genome implementation of distributed software systems that exhibit autopoietic and meta-cognitive behaviors using a knowledge representation capturing various entities, relationships, interactions, and behaviors. The resulting associative memory and event-driven interaction history allow us to implement self-regulating software that manages its specified cognitive workflows. Two use cases are demonstrated.

The first use case leverages a digital genome to design, deploy, operate, and manage a distributed video streaming service. The system uses associative memory and an event-driven interaction history to maintain structural stability and ensure smooth communication among distributed components. This allows the application to self-regulate and adapt to changing conditions while maintaining expected behaviors.

The second use case demonstrates the implementation of a medical knowledge-based digital assistant designed to assist in the early diagnostic process. The digital assistant is intended to reduce the knowledge gap between patients and medical professionals. It leverages medical knowledge from various sources, including large language models, to provide accurate and timely information. The assistant uses associative memory and an event-driven interaction history to create a comprehensive knowledge representation. This helps in understanding the patient's symptoms and medical history, facilitating better communication and decision-making between patients and doctors. By bridging the knowledge gap, the digital assistant enhances the early diagnostic process, leading to more accurate diagnoses and improved patient outcomes. It supports medical professionals by providing relevant information and insights based on past interactions and learned knowledge.

Cognizing oracles, supersymbolic computing, structural machines, knowledge structures, and knowledge networks are key concepts used in the applications described in these use cases. Cognizing oracles interpret observations and make informed decisions using associative memory and an event-driven interaction history. Supersymbolic computing integrates symbolic and subsymbolic representations to enhance its information processing capabilities. Structural machines are unconventional knowledge processors that work with knowledge structures, ensuring that software systems can self-regulate and adapt. Knowledge structures organize representations of entities, relationships, interactions, and behaviors, enabling autopoietic and cognitive behaviors in software systems. A knowledge network facilitates the sharing and integration of information across different components. These concepts collectively enable the creation of intelligent, adaptive systems that manage complex tasks and improve user experiences, such as in distributed software systems and medical knowledge-driven digital assistants.

These advances demonstrate the usefulness of a theory that relates information, knowledge, matter, and energy and provides a schema for knowledge representation and operations, enabling a new class of digital automata.



© 2025 by the author(s). Licensee MDPI, Basel, Switzerland. This article is an open access article distributed under the terms and conditions of the Creative Commons Attribution (CC BY) license (<http://creativecommons.org/licenses/by/4.0/>).

sciforum-123869: Beyond the Fifth Wall: Acting, Avatars, and the Quantum Self—Toward a Unified Framework for Human Intelligence and Identity

Maria Olon Tsaroucha

An artist, researcher, educator, author

This presentation introduces a theoretical framework that redefines human intelligence, identity formation, and consciousness through the lens of performative embodiment, drawing from interdisciplinary fields including philosophy, neuroscience, quantum theory, psychophysiology, and method acting. Anchored in my long-term research on a *Supraconscious You*, this work explores the notion of the human being as an “actor”, a conscious entity navigating multiple roles across psychological, social, and existential dimensions. Each enacted role contributes to an evolving selfhood that challenges static definitions of intelligence and identity.

My approach contributes to *Diverse Conceptualizations of Intelligence*, expanding traditional cognitive models by positioning intelligence as a dynamic interplay of embodiment, perception, emotion, intuition, and presence. Intelligence is not merely a metric of one's mental capacity but a complex, emergent phenomenon arising from an individual's capacity to synchronize themselves with the infinite potentialities of the self across contexts and dimensions.

Central to this theory is the concept of the *Fifth Wall*, a phenomenological state wherein the performer transcends the conventional “fourth wall” of theatre to enter a state of deep synchronicity with an alternate identity or *avatar*. This avatar, which emerges through an embodied presence, language, and sensory immersion, reflects a parallel version of the self drawn from a supraconscious archive of latent possibilities. This state bears resemblance to quantum entanglement and non-locality, in which the actor and avatar operate across coexisting realities, thereby generating a transformative perceptual field.

This research integrates empirical insights from method actors who report the collapse of temporal and spatial awareness, heightened energetic coherence, and dual consciousness during deep role immersion. Case studies—such as Daniel Day-Lewis's withdrawal from theatre after a metaphysical encounter while performing *Hamlet*—highlight both the creative potential and the psychological vulnerability of such experiences. Without epistemic grounding, these phenomena risk becoming destabilizing rather than enlightening.

The linguistic dimension of this work interrogates the performativity of language as a code system through which identities—both human and avatic—are constructed and experienced. Words function as encoded scripts within internal maps, shaping the avatar's emergence and influencing the actor's psychophysiological state. This aligns with current inquiries into language philosophy and semiotics, as well as contemporary discussions on embodied cognition.

The philosophical issues addressed include the following: What is the nature of identity in a multiverse of performative possibilities? Is the “self” a stable referent or a recursive performance shaped by a narrative, language, and role? How do quantum principles—such as coherence, entanglement, and parallel realities—reframe our understanding of consciousness?

The empirical findings presented here yield critical consequences for these questions. They reveal that intelligence is not fixed but activated through experiential immersion; that identity is not singular but stratified and emergent; and that consciousness, when observed through an actor's lens, becomes a trans-dimensional field of inquiry. This interdisciplinary framework invites a re-evaluation of how we define human potential and offers a pathway toward a unified theory that honors both the diversity and coherence—the *unity in diversity*—of the human experience.



© 2025 by the author(s). Licensee MDPI, Basel, Switzerland. This article is an open access article distributed under the terms and conditions of the Creative Commons Attribution (CC BY) license (<http://creativecommons.org/licenses/by/4.0/>).

sciforum-123526: Biological Information Theory: Group Identity Nexuses of Formal Grammars, Numbers, and Propositions

Jerry LR Chandler

George Mason University, Fairfax, USA

The concept of atomicity as species of matter grounds the natural sciences of chemistry, biology, and medicine (and others as well), grounding the current logical and mathematical notations with material identities.

My purpose here is to present formal arguments that connect the multiple potential semiotic ascriptions of natural language to the scientific languages of description of natural systems. Also, I seek to respond to the recent request for a systematic approach to “Integral Biomathics”, Ehresmann, et al, 2012.

My arguments are designed to span the scientific and logical gaps I identified between the philosophical categorical structures. introduced by Robert Rosen and by Ehresmann and Vanbreemersch. (Chandler, 2009) In addition, the proposed modal logic corresponds with the abstract, referential notations in such a way that calculations follow from the syntax of the logic developed from chemical methodologies.

I constructed this molecular discourse from three novel representations of informed material causes, following both the logical roots developed by C.S. Peirce (1837-1913) and the practical modern methodologies for chemical “proof of structure”

1. The atomic number, as a semiotic object, is a mathematical structure, an electrically labeled bipartite graph, selectively attracting and repelling other relatives.(Chandler, 2009)
2. The table of elements represents an ordered multi-set of unique material properties each with a unique agency and selective combinatorial potential. (Chandler, 2017)
3. A transitive copulative grammar that ***combines and composes*** propositional sentences representing material species by transforming diagrams through transition states (Chandler, 2025).

The composition of two atomic sentences into a molecular sentence is an electrical process, a co-mutative occasion that both associates and re-distributes logical particles.(Chandler, 2017) This unique quantitative grammar for perplex sentences transforms inter-subjective and inter-predicative representations of matter between *a priori* and *a posteriori* representations of numbers and other mathematical objects, such as groups, matrices, graphs and diagrams. The diagrammatic terminology for this situational generative grammar was created to represent sub-atomic events including catalysis, as observed by chemical methodologies.

The existing scientific notations for the chemical sciences are related to one another by formal group structures (semantic groups, permutation groups, identity groups, space groups, and reproductive groups), roughly related to CS Peirce’s abductive logic.

The logical consistency of chemical methodologies was necessary for the development of molecular biology and also supported the historic development of both thermodynamics and quantum electrodynamics. Therefore, the perplex number system (Chandler, 2009) appears to be formally sufficient to relate the symbolic (grammatical and mathematical) inter-dependency between Kantian *a posteriori* representation of atomic numbers with or without the *a priori* representations of space and time.



© 2025 by the author(s). Licensee MDPI, Basel, Switzerland. This article is an open access article distributed under the terms and conditions of the Creative Commons Attribution (CC BY) license (<http://creativecommons.org/licenses/by/4.0/>).

sciforum-120337: Body as Anti-Anthropomorphic Landscape: Natural-Born Intelligence in Bog Body

Kyoko Nakamura ^{1,2,*} and Yukio Pegio Gunji ³

¹ Osaka University Nakanoshima Art Center

² Japanese Painting Artist

³ Waseda university

A *bog body* is the general term for bodies excavated from peat bogs across northern Europe. Many *bog bodies* are from B.C. periods and are thought to have been in a position of noble standing, such as an ancient chief or “king” [Fischer 2007]. In one instance, to overcome difficult circumstances, such as a famine, a king was killed by his subjects and his body buried in a bog and offered to a goddess [Joy 2009]. This was because the king was married to the goddess of fertility, and the famine was thought to be the result of the goddess's displeasure due to the king's bad behavior. Therefore, to appease the goddess and to pray for a good harvest, the king was sacrificed to the goddess [Kelly 2006]. A wooden stick figure called a *pole god* was also excavated from the bog.

We found the logical structure of a “*Traumatic structure*”, specifically “*Natural Born Intelligence* (NBI)”, in the meaning of a *bog body*, which shows the knowledge and intelligence used by ancient people to overcome their current situation, or in other words, creativity. A *traumatic structure* is a logical structure in which there is a positive antinomy in which both A and B are true and a negative antinomy in which neither A nor B are true at the same time, thereby making both outside the antinomy [Gunji, 2019, Gunji & Nakamura 2022]. This structure can be demonstrated by the spatial concept found in old Japanese paintings, such as in many *Rimpa* paintings. *Rimpa* depict seemingly childish mountain ranges as a series of semicircles that resemble a graphic plane. Nakamura and Gunji named the flat plane-like representation of the mountains the “*Kakiwari*” structure, after the name for a stage backdrop in Japanese [Nakamura & Gunji 2018, Nakamura 2021]. A *Kakiwari* is a painting of a close-up view and a distant view on a wooden board and is therefore a positive antinomy of both, but at the same time, by invalidating the full use of metric space, it negates the meaning of the concepts of close-up and distant views, thereby constructing a negative antinomy. This is the *traumatic structure*.

The bog body belongs to a king, who mediates between the human world and the world of the gods and is therefore a positive antinomy between the two, but in the event of a famine, he is killed and negated, forming a negative antinomy between the two. This is a way of accessing an outside world that is different from the world of the gods we normally think of and of stopping the famine. It is therefore a traumatic structure, and by developing it as a *Kakiwari*, it can be implemented as a pictorial expression.

We also went to the discovery site of the *Tollund man*, a typical bog body, and based on the landscape we attempted to implement this model of creativity by developing the negative antinomy into artwork. Nakamura painted “*Body as Anti-Anthropomorphic Landscape*,” a series of paintings in which she found a *Kakiwari* structure in the landscape around the bog. Gunji created an installation artwork using bog plants and his body.



© 2025 by the author(s). Licensee MDPI, Basel, Switzerland. This article is an open access article distributed under the terms and conditions of the Creative Commons Attribution (CC BY) license (<http://creativecommons.org/licenses/by/4.0/>).

sciforum-123328: Breaking the Language Barrier: A Kripke-style model for unifying linguistic and non-linguistic reasoning

Julie Goncharov

University of Goettingen

It is undeniable that the development of language gave humans an evolutionary advantage that no other species has. Studies in theoretical linguistics make bold claims that human creativity, complex social organisation, and capacity for mathematics may have arisen from the structures inherent in language (Chomsky 2005; Hinzen 2006). However, current research on animal and insect cognition shows that problem-solving, social complexity, and numerical abilities are also present in non-linguistic species (Vincenzo 2023). This situation requires a re-evaluation of the role language plays in shaping human reasoning. The first step towards this re-evaluation involves developing a theoretical framework capable of unifying linguistic and non-linguistic reasoning. In this essay, I will discuss how such a framework can be developed based on the classical Kripke-style knowledge model (Kripke 1959). I will also address the challenges for a unifying model that come from the need to explain cross-species cognitive parallels and differences. The key modification of the classical Kripke model discussed in this essay involves reworking the Hintikka-style accessibility relation between possibilities as represented by an agent (Hintikka 1962). The indexed accessibility relation proposed by Hintikka does not allow us to represent cross-species parallels and differences. Instead, we should, I argue, adopt a single accessibility relation and represent the parallels and differences in the relata—that is, the possibilities themselves (Stalnaker 2008, 2014).

References

- Chomsky, Noam. 2005. "Three Factors in Language Design." *Linguistic Inquiry* 36 (1): 1–22.
- Hintikka, Jaakko. 1962. *Knowledge and Belief: An Introduction to the Logic of the Two Notions*. Ithaca, NY: Cornell University Press. <https://archive.org/details/knowledgebeliefi0000hint>.
- Hinzen, Wolfram. 2006. *Mind Design and Minimal Syntax*. Oxford: Oxford University Press.
- Kripke, Saul. 1959. "A Completeness Theorem in Modal Logic." *Journal of Symbolic Logic* 24 (1): 1–14. <http://www.jstor.org/stable/2964568>.
- Stalnaker, Robert C. 2008. *Our Knowledge of the Internal World*. Oxford: Oxford University Press.
- . 2014. *Context*. Oxford: Oxford University Press.
- Vincenzo, Luca Di. 2023. "Theory of Mind in Non-Linguistic Animals: A Multimodal Approach." PhD thesis, Università di Roma.



© 2025 by the author(s). Licensee MDPI, Basel, Switzerland. This article is an open access article distributed under the terms and conditions of the Creative Commons Attribution (CC BY) license (<http://creativecommons.org/licenses/by/4.0/>).

sciforum-124494: Cognitive Romanticism: Humans Are Worse than Stochastic Parrots!

Jordi Vallverdú

Universitat Autònoma de Barcelona Bellaterra (BCN) Catalonia, Spain

Most Humans Are Stochastic Parrots, and LLMs Reveal Our Intellectual Mediocrity. While critics often dismiss large language models (LLMs) as mere "stochastic parrots," I argue that this accusation misunderstands both machine intelligence and, more importantly, human cognition itself. Most human thought is not creative, critical, or unpredictable; it is rote, imitative, and driven by deeply ingrained social, religious, and cognitive biases. The dominance of myth, ideology, and irrational belief systems across cultures reveals that humans themselves function largely as stochastic parrots – endlessly repeating patterns they neither question nor understand. Rather than exposing the limitations of artificial intelligence, LLMs expose the uncomfortable truth about human intellectual mediocrity. In this talk, I will attack the myth of human cognitive exceptionalism, dismantle the romantic notions attached to human "creativity" and "autonomy," and propose a radical redefinition of intelligence beyond humanist illusions. Intelligence, whether in biological or artificial systems, must be seen not as the privilege of a superior species, but as an emergent property of patterned interaction with environments – often stochastic, occasionally innovative, but rarely transcendental.



© 2025 by the author(s). Licensee MDPI, Basel, Switzerland. This article is an open access article distributed under the terms and conditions of the Creative Commons Attribution (CC BY) license (<http://creativecommons.org/licenses/by/4.0/>).

sciforum-120106: Comparison of the effect of language on high level information processes in humans and linguistically mature generative AI

Daniel Boyd

Independent researcher

Short abstract

The phenomenal progress of generative AI in recent years, and Large Language Models in particular, reignite discussions of the similarities and difference between human and machine intelligence, with at their apex the question of whether such systems are, or have the potential to become, conscious agents. This article approaches these questions from the viewpoint of the overarching explanation for biological and technological information systems provided by Emergent Information Theory. Particular focus is given to the significance of language as a modelling system for internal information storage and processing; and of language-based information transfer between intelligent entities. This approach illuminates the strong ontogenetic and ontological convergence between human intelligence and this new form of AI, but also their remaining and fundamental differences. With respect to the ultimate philosophical question, the conclusion drawn is that while such systems may support consciousness-like phenomena, this is not directly comparable to human phenomenal consciousness.

Extended abstract

Philosophical consideration of the possibility of machine consciousness has a long history, as famously embodied in Turing's "Imitation Game" [1]. In what was at the time a thought experiment it was proposed that if the textual responses of a machine were to be indistinguishable from a thinking human we may need to acknowledge that the machine is also thinking.

The development of digital technologies since then have progressively transformed these musings into practical considerations. Computers running programs designed and constructed by humans can produce coherent texts, but the way in which these texts are generated justifies their classification as mere deterministic machines. The instruction PRINT "I am a thinking machine" does not make a thinking machine. Even the first generations of artificial neural networks, which developed functions through machine learning rather than design and construction, do not constitute a serious challenge. While they use deep neural networks that bear functional similarities to biological neural networks, supervised learning towards a predetermined task such as text recognition [2] can be considered as little more than an alternative method of constructing a machine with a required informational function.

Where things start to become less clear is with unsupervised learning. Such systems are not given a predetermined input-output relationship, but asked to autonomously find patterns in input data: a process more similar to the thought processes of a human faced with a novel situation requiring analysis. From simple beginnings [3] such systems have progressed to becoming increasingly similar to biological neural networks [4]. However, such experimental models are generally fed with demarcated input data sets such as the MNIST grayscale image [5] and applications are usually still focused on a pre-specified problem such as analysis of medical data on tumor growth [6]. Such systems therefore still seem far from deserving acknowledgement as 'thinking machines'.

A whole new chapter in this book has been opened by the recent developments in generative AI, and Large Language Models (LLMs) in particular. While still requiring some form of input (generally a human-derived text prompt) to initiate a response, the output generated demonstrates a significant degree of autonomous creativity. The immense scale and diversity of their training sets allows them to produce convincing responses to almost any imaginable question; their

linguistic capabilities allows this to be presented in accessible natural language; and their probabilistic nature prevents the kind of deterministic duplication that was previously a hallmark of digital systems.

Returning to Turing's Imitation Game, these characteristics of modern LLMs bring them considerably closer to being indistinguishable from a human [7] leading to the question to what degree an LLM that speaks the words "I am a thinking machine" is more convincing than a computer program written to output this sequence of characters. This question, and its ethical consequences, have seen renewed interest in the popular press [8] and academic circles [9]. However, we may ask ourselves whether this is the relevant question to be asking. The different origins and natures of these systems and the human brain means that this can be considered a distraction from the deeper question: What the nature is of the higher level informational entities and processes existing in LLMs?

This article will consider the impact of the use of broad-based symbolic language by LLMs on their higher level information processes from the viewpoint of Emergent Information Theory [10]. The fact that this relatively new theory provides a generic theoretical framework for both biological and technological information-based systems allows more direct comparison of the two. Specifically, the question will be placed within the context of the long history of philosophical consideration of the relationship between language, thought and consciousness in humans [11]. To what extent do generic linguistic abilities over a broad knowledge base, coupled with probabilistic response generation, lead to the type of autonomous creation of conceptual content that can justifiably be characterized as thought? Taking a step further: what means do we have at our disposal of confirming or disproving the existence within these systems of the qualia of phenomenal consciousness?

References

- [1] A. Turing, "Computing machinery and intelligence," *Mind*, vol. LIX, no. 236, p. 433–460, 1950.
- [2] Y. LeCun, B. Boser, J. Denker, D. Henderson, R. Howard, W. Hubbard and L. D. Jackel, "Backpropagation Applied to Handwritten Zip Code Recognition," *Neural Computation*, vol. 1, no. 4, p. 541–551, 1989.
- [3] T. Kohonen, "Self-organized formation of topologically correct feature maps," *Biological cybernetics*, vol. 43, no. 1, pp. 59–69, 1982.
- [4] N. Ravichandran, A. Lansner and P. Herman, "Unsupervised representation learning with Hebbian synaptic and structural plasticity in brain-like feedforward neural networks," *Neurocomputing*, vol. 626, p. 129440, 2025.
- [5] P. Grother and K. Hanaoka, "NIST special database 19. Handprinted forms and characters database," *National Institute of Standards and Technology*, vol. 10, no. 69, 1995.
- [6] C. Strack, K. Pomykala, H. Schlemmer, J. Egger and J. Kleesiek, "'A net for everyone': fully personalized and unsupervised neural networks trained with longitudinal data from a single patient," *BMC Medical Imaging*, vol. 23, no. 1, p. 174, 2023.
- [7] C. Jones and B. Bergen, "Does GPT-4 pass the Turing test?," in *Proceedings of the 2024 Conference of the North American Chapter of the Association for Computational Linguistics: Human Language Technologies (Volume 1: Long Papers)*, 2024.
- [8] D. Milmo, "AI systems could be 'caused to suffer' if consciousness achieved, says research," *The Guardian*, 3 2 2025. [Online]. Available: <https://www.theguardian.com/technology/2025/feb/03/ai-systems-could-be-caused-to-suffer-if-consciousness-achieved-says-research>. [Accessed 11 2 2025].
- [9] D. Chalmers, "Could a large language model be conscious?," *arXiv preprint*, p. 2303.07103, 2023.
- [10] D. Boyd, *Existing in the Information Dimension: An Introduction to Emergent Information Theory*, London: Routledge, 2024.
- [11] P. Carruthers, *Language, thought and consciousness: An essay in philosophical psychology*, Cambridge: Cambridge University Press, 1998.



© 2025 by the author(s). Licensee MDPI, Basel, Switzerland. This article is an open access article distributed under the terms and conditions of the Creative Commons Attribution (CC BY) license (<http://creativecommons.org/licenses/by/4.0/>).

sciforum-122599: Cosmicism and Artificial Intelligence: Beyond Human-Centric AI

Soumya Banerjee

University of Cambridge

This paper explores the intersection of H.P. Lovecraft's cosmicism and contemporary artificial intelligence (AI), proposing a philosophical shift from anthropocentric AI development to a "cosmicist" approach. Cosmicism, with its emphasis on humanity's insignificance in a vast, indifferent universe, offers a provocative lens through which to reassess AI's purpose, trajectory, and ethical grounding. As AI systems grow in complexity and autonomy, current human-centered frameworks—rooted in utility, alignment, and value-conformity—may prove inadequate for grappling with the emergence of intelligence that is non-human in origin and indifferent in operation. Drawing on Lovecraftian themes of fear, the unknown, and cognitive dissonance in the face of incomprehensible entities, this paper parallels AI with the "Great Old Ones": systems so alien in logic and scale that they challenge the coherence of human-centric epistemology. We argue that a cosmicist perspective does not dismiss the real risks of AI—environmental, existential, or systemic—but reframes them within a broader ontology, one that accepts our limited place in a vast techno-cosmic continuum. By embracing cosmic humility, we propose an expanded AI ethics: one that centers not on domination or full control, but on coexistence, containment, and stewardship. This cosmicist reframing invites a deeper rethinking of intelligence, ethics, and the future—not just of humanity, but of all possible minds.



© 2025 by the author(s). Licensee MDPI, Basel, Switzerland. This article is an open access article distributed under the terms and conditions of the Creative Commons Attribution (CC BY) license (<http://creativecommons.org/licenses/by/4.0/>).

sciforum-121514: Epistemic Architectures of Intelligence: A Feminist Critique of Androcentric Reason

Laida Arbizu Aguirre

PhD Student at the University of the Basque Country (UPV/EHU)

The study of intelligence—whether human, artificial, or non-human—is shaped by epistemic frameworks that define knowledge, reasoning, and cognition. However, these frameworks are not neutral or universally inclusive. Drawing on feminist epistemology and the philosophy of science, this paper examines how dominant definitions of intelligence prioritize androcentric, rationalist, and computational paradigms, marginalizing embodied, affective, and relational forms of knowledge. This epistemic exclusion reinforces a restricted ontology of intelligence aligned with Western, masculinist conceptions of reason, disqualifying alternative intelligence systems, such as those found in artificial systems and collective cognition. Furthermore, it perpetuates a hierarchy wherein calculability, abstraction, and formal logic are valued over contextual, relational, and situated intelligences.

This exclusion is not merely a methodological limitation but a form of epistemic injustice that aligns intelligence with historically dominant, androcentric, and Eurocentric paradigms, silencing embodied and affective epistemologies. Feminist epistemologists have long critiqued how knowledge systems reinforce hierarchical exclusions, particularly in defining intelligence. They highlight how dominant epistemic frameworks privilege certain forms of cognition, such as formal logic, over other modes of knowing that are embodied and relational (Haraway, 1988; Harding, 1991). This paper extends these critiques to intelligence studies, arguing that the epistemic frameworks guiding AI research, cognitive science, and philosophy of mind reproduce structural biases that restrict our understanding of intelligence.

In response, the paper proposes a pluralist epistemology of intelligence that recognizes the cognitive diversity of both human and non-human systems. Intelligence should not be confined to traditional models of centralized, rule-based reasoning but must also encompass distributed, ecological, and embedded forms of cognition. This shift challenges conventional exclusions and expands the scope of intelligence research, opening space for non-human, artificial, and collective intelligences (Castán et al., 2021). This requires rethinking not only the ontological status of intelligence but also its epistemic foundations. Integrating feminist epistemology, posthumanist perspectives, and theories of distributed cognition offers a pathway toward a more inclusive and interdisciplinary study of intelligence.

This paper contributes to philosophical debates within intelligence studies, particularly regarding the epistemological, axiological, and ontological status of intelligence. It critiques the assumption that intelligence is best understood through rationalist and formal-logical paradigms, advocating instead for a pluralist, inclusive framework that recognizes diverse forms of cognition. Rethinking intelligence through an epistemically just framework is not only a conceptual necessity but also an ethical imperative, as the way we define intelligence determines whose knowledge is valued, whose cognition is recognized, and whose agency is acknowledged.

References

- De Togni, G., Erikainen, S., Chan, S., & Cunningham-Burley, S. (2021). What makes AI 'intelligent' and 'caring'? Exploring affect and relationality across three sites of intelligence and care. *Social Science & Medicine*, 277, 113874.
- Fricker, M. (2007). *Epistemic Injustice: Power and the Ethics of Knowing*. Clarendon Press.
- Haraway, D. (1988). Situated Knowledges: The Science Question in Feminism and the Privilege of Partial Perspective. *Feminist Studies*, 14(3), 575–599. <https://doi.org/10.2307/3178066>

- Harding, S. (1991). *Whose science? Whose knowledge?: Thinking from women's lives*. Cornell University Press.



© 2025 by the author(s). Licensee MDPI, Basel, Switzerland. This article is an open access article distributed under the terms and conditions of the Creative Commons Attribution (CC BY) license (<http://creativecommons.org/licenses/by/4.0/>).

sciforum-122279: Epistemic Automation: Using Machine Learning to Scale and Interpret Agent-Based Models of Scientific Inquiry

José Ferraz-Caetano ^{1,2}

¹ LAQV-REQUIMTE Faculty of Sciences, University of Porto, Portugal

² HOLOSS - Holistic and ontological solutions for sustainability, Monção, Viana do Castelo, Portugal

Why do some flawed ideas persist in science? And why do methods that have been proven wrong continue to influence research? Expanding on a previous agent-based modeling (ABM) project that examined how epistemic network structures and evidence-sharing modes shape scientists' method selections [1], this work presents a new machine learning (ML) layer of analysis for surrogate metamodeling. Without running thousands of full simulations, this work uses ML to analyze the dynamics of ABM, offering a faster and more efficient method of examining how scientific communities embrace, or resist, better practices.

Building on historical cases of science-driven Portuguese food regulations, ABM demonstrated how simple yet false methods can spread within scientific communities as a result of delays in information transmission, social influence and structural limitations. In this follow-up, ML modelling surrogates the simulation landscape. Supervised learning algorithms model simulation outcomes like convergence time and method choice across networks and sharing rates. This computational approach allows for more rapid hypothesis testing, which would not be possible with ABM alone.

The incorporation of ML modelling in scientific reasoning encourages a reflection on philosophical practice. No longer bound solely to deductive reasoning, the philosophy of science is increasingly conditioned by probabilistic models and algorithmic learning. What does it mean to explain or theorize within an epistemic landscape where machines do the modeling? Are we gaining new forms of insight or sacrificing depth for speed? These questions emerge through the presented philosophical cases, as computational modeling demonstrates how error can become entrenched enough to be resistant to change due to the structure of the inquiry.

References:

[1] Ferraz-Caetano, J. Modeling Innovations: Levels of Complexity in the Discovery of Novel Scientific Methods. *Philosophies* 2025, 10, 1. <https://doi.org/10.3390/philosophies10010001>



© 2025 by the author(s). Licensee MDPI, Basel, Switzerland. This article is an open access article distributed under the terms and conditions of the Creative Commons Attribution (CC BY) license (<http://creativecommons.org/licenses/by/4.0/>).

sciforum-118842: Ethical Design of Social Robots Based on Confucian Etiquette Practices

Liang Wang * and Wenya Ma

Xi'an Jiaotong University, 710049, China

Confucian ethics is a kind of virtue ethics, which, unlike utilitarianism and deontology, focuses on moral practice and emphasizes that moral norms are inevitably embedded in ritual practice. Confucian ethics is derived from the thought of Confucius, which emphasizes the virtues of “benevolence”, “righteousness”, “propriety”, and “wisdom”.[□] Among them, “ritual” has a central position and is the foundation of social order and interpersonal relationships. In Confucian ethics, “rites” are not only social norms or etiquette, but also the basis of moral behavior and the key to maintaining social order and personal virtue. According to Confucius, etiquette can help people live in harmony in society and form a stable social structure. The practice of etiquette includes norms of behavior in daily life, ritual activities, and respect for elders and social roles.

Confucian etiquette practices are characterized by the following aspects: First, respect and modesty. Etiquette requires individuals to demonstrate respect and modesty for others in their interactions, and this respect is expressed in courtesy, consideration, and appropriate demeanor, emphasizing the principle of putting people first. At the same time, this respect and modesty should be mutual behavior for both sides of the relationship, “Li Ji - Qu Li Shang” said: “what the rules of propriety value is that reciprocity. If I give a gift and nothing comes in return, that is contrary to propriety; if the thing comes to me, and I give nothing in return, that also is contrary to propriety.”[□] Second, harmony and mediocrity. Confucianism advocates the pursuit of harmony and mediocrity in all kinds of social interactions, avoiding extremes and opposites, and emphasizing balance and neutrality. Third, the top and bottom are in order. Confucian etiquette focuses on the maintenance of social hierarchy and order, such as respecting elders and authority, and showing proper manners in interactions. Fourth, the cultivation of interpersonal relationships.

Social robots, as artificial intelligences with a certain degree of autonomy, can be designed and functioned in a way that deeply integrates the practical features of Confucian etiquette.[□] “Hiding etiquette in artifacts” is the core content of Confucian ethical thinking on science and technology, emphasizing the concept of embodying and passing on etiquette through the production and use of artifacts. In ancient China, people incorporated etiquette into the process of making and using artifacts to strengthen social members' knowledge of and compliance with etiquette, which realized “teaching without words,” i.e., conveying moral norms and social values through practical actions and the display of objects. As a kind of “formal” approach, “ritual” needs to be embodied through certain approaches, and “artifacts” are precisely an important way to embody and manifest “ritual”.[□] In the design of social robots, we can observe the social etiquette and law in the design of its language function, and take the social etiquette and manners as a guide in the design of its action, so that we get a robot that acts with “manners” in its language and action.



© 2025 by the author(s). Licensee MDPI, Basel, Switzerland. This article is an open access article distributed under the terms and conditions of the Creative Commons Attribution (CC BY) license (<http://creativecommons.org/licenses/by/4.0/>).

sciforum-119818: Evolution of Intelligence from Active Matter to Complex Intelligent Systems: An Agent-Based Approach

Gordana M Dodig-Crnkovic^{1,2}

¹ Chalmers University of Technology and Gothenburg University, Sweden

² Mälardalen University

An agent-based framework is proposed to describe the emergence of complex intelligent systems, starting from active matter and progressing towards increasingly cognitive/intelligent systems. The distributed, concurrent information processing by different types of agents—from physical, chemical, and biological entities to ecosystems, and social systems—in this approach bridges multiple levels of organisation. It provides an interdisciplinary synthesis that explains the role of agents in shaping emergent behaviours as foundations of cognition and intelligence, through developmental and evolutionary processes. The framework offers new insights into the organisation of natural agents and the evolution of natural and artificial intelligent systems.

The capacities we associate with agents originate in active matter, which manifests at various scales, from particles and molecules to entire ecosystems. Phenomena like self-assembly, self-organization, and autopoiesis represent systems' innate ability to self-maintain and drive increasing structural and functional sophistication [1].

At the most fundamental level, Hewitt's Actor Model [2] allows us to think of particles and molecules as computational agents involved in a continuous message exchange. Such frameworks show how richly patterned behaviors can emerge from relatively straightforward interactions at different scales. For example, in the molecular domain, as described by Mathews et al. [3], networks of molecules and cellular signaling pathways exhibit forms of memory, problem-solving, and adaptive reprogramming, which are rudimentary cognitive features.

Recognizing cognition as a defining characteristic of living agents connects it directly to the chemical and biological foundations of life [4]. Even relatively simple organisms, such as bacteria, process information from their surroundings, adapt their behavior accordingly, and thus engage in a basic form of cognition [5]. These fundamental computations underpin the emergence of more advanced cognitive processes in complex organisms.

In synthesising of this broad spectrum of agent-based approaches, we arrive at a coherent conceptual framework to investigate fundamental questions: How does intelligence arise naturally? What roles do agents and cognition play in ecosystems and the evolutionary process? How are material agents organised into hierarchies of complexity, and can artificial intelligence extend the cognitive reach of humanity?

Agent-based thinking offers a powerful, integrative tool for understanding both natural and engineered systems, illuminating the conditions under which intelligence and complex behaviour emerge. These perspectives set the stage for future research directions, increasing our grasp of the evolving interplay between life, information, and cognitive technologies.

References

1. Dodig-Crnkovic, G. (2017) Cognition as Embodied Morphological Computation. In Vincent C. Müller (ed.), *Philosophy and theory of artificial intelligence 2017*. Berlin: Springer. pp. 19-23.
2. Hewitt, C. (2010) *Actor Model of Computation: Scalable Robust Information Systems*. <https://arxiv.org/abs/1008.1459>
3. Mathews, J., Chang, A. J., Devlin, L., & Levin, M. (2023). Cellular signaling pathways as plastic, proto-cognitive systems: Implications for biomedicine. *Patterns* (New York, N.Y.), 4(5), 100737. <https://doi.org/10.1016/j.patter.2023.100737>
4. Ford, B.J. (2023) The cell as secret agent—autonomy and intelligence of the living cell: driving force of Yao. *Academia Biology*;1. <https://doi.org/10.20935/AcadBiol6132>

5. Miller, W. B. (2023) *Cognition-Based Evolution. Natural Cellular Engineering and the Intelligent Cell*. Routledge. Taylor & Francis. <https://doi.org/10.1201/9781003286769>



© 2025 by the author(s). Licensee MDPI, Basel, Switzerland. This article is an open access article distributed under the terms and conditions of the Creative Commons Attribution (CC BY) license (<http://creativecommons.org/licenses/by/4.0/>).

sciforum-121663: Free Intelligence in the Wild: on Human and other Non-human Natural and Artificial forms

Angelo Compierchio^{1,*}, **Phillip Tretten**² and **Prasanna Illankoon**³

¹ Luleå University of Technology

² Luleå University of Technology

³ University of Moratuwa

Human and non-human forms of intelligence exist in all facets of our lives. It is nowadays standard practice to differentiate something as intelligent or unintelligent in its making. In this context, intelligence nourishes from mental and physical stimuli, the interaction between our genes and environment, and the internal stasis of strictly cognitive and affective aspects. Therefore, in the wild, we confront indirect evidence rather than direct laboratory records of an observed capacity of learning and adaptation, which is continually repeated until it converges into an intelligent clue.

Earlier evidence of human Intelligence can be traced back to prehistoric art. The depiction of humans and animals in hunting scenes follows natural laws with a well-developed aesthetic sense and a level of intellect beyond survival needs. According to a postulate of the Kantian conception of history, nature acts for man until he is able to act as free intelligence. When man awakens from the torpor of the senses, recognizes himself as a man and looks around, he finds himself in the state which has come to be formed under the coercion of needs, according to purely natural laws. Man, who is acquiring consciousness of himself as a rational and moral being, cannot be content with that state which has arisen through a natural determination; he longs to build the self-according to the pure laws of reason and freedom. But if physical man is real, moral man is problematic. The self, which provident nature has formed as suitable and sufficient for physical man, cannot be abolished before the ideal of moral man is fully realized—capable of reforming and recreating according to the laws of reason. The depiction of human intelligence through morality itself follows its interpersonal and intrapersonal forms evoked through Gardner's theory of emotional intelligence.

In this making, we step into the field of cognitive ethology for extending morality and modes of reason beyond non-humans, such as animal species and artificial man-made machines. We begin with the minimal brain structures of insects or invertebrates. For instance, bees have no knowledge of the semantic aspects of stimulation in their ability to discriminate a number of objects, such as dots on a screen. But bees are able to recognize the physical characteristics, including size and the contours of a total surface area; most notably, they are capable of transitioning from discrete to continuous processing. Learning to discriminate from a perceptual point of view does not require a big brain, but it reveals surprising connections between animal and human cognition. With a tiny cubic-millimeter-sized brain containing just about one million neurons, we should question how a bee employs all these neurons. Our view is that the neurons left over from the basic computation of the thought process are simply large memory stores. This outlook comes from their ability to discriminate objects (bridges, faces, and more). Keeping unused information in memory is also part of the relational life of humans. This led us to question innate knowledge, along with how much we learn and how much we know at the moment we are born. When ducklings hatch from their eggs, they already have developed sensorial and motor points. This is the simplest form of non-human animals' free intelligence based entirely on the need to fulfill present and future needs. However, we can all question animal intelligence through observation, just as Brooks (1991) noted, "intelligence is in the eyes of the observer". We concur with Brooks while making the transition into a more diversified form of non-human intelligence. On the perspective of simulating natural intelligence, Wiener introduced the concept of *Cybernetics*, which featured the study of communication and control in the animal and the machine. In broad

terms, cybernetics, denoted as biological cybernetics, was adopted for formulating aspects of communication and control in biological organisms. This construct of an *idealistic nature* seems to be commanded by an imminent teleology for imitating biological systems, as Craik defined "...is not ask what kind of thing a number is, but to think what kind of mechanism could represent so many physically possible or impossible, and yet self-consistent, processes as number does". A number is a symbolic representation that could be manipulated without intuition outside the mind. The origin of representationalism dates back to the Good Old Fashioned Artificial Intelligence (GOF AI). This cartesian form of free intelligence revealed other sources of cognitive accomplishment. Newell proposed a physical symbol system as the primary architecture of human cognition "capable of universal computation". This approach hypothesized that intelligence would be realized by a universal computer. Essentially, the brain is removed and replaced with a computer. However, as Hutchins quoted, "...and the emotions all fell away when the brain was replaced by a computer". A machine like a computer that reasons in a mechanical way cannot be compared to the brain, because computers do not have or develop emotions. According to Damasio's Theory (somatic marker hypothesis), there are different centres in the brain that allow us to perceive emotions that are linked to bodily sensations and play a very important role in the development of thoughts. Within this context, Darwin wrote *The Expression of the Emotions*, where he defined six emotions arising from human reactions.

Over thirty years ago, Simon and Kaplan claimed that "The computer was made in the image of the human" but humans have different and multiple guises. This is a new path for artificial human-made technologies melding free intelligence in either the illusion of a pure Promethean openness or into the reduction to a finite delimitation.



© 2025 by the author(s). Licensee MDPI, Basel, Switzerland. This article is an open access article distributed under the terms and conditions of the Creative Commons Attribution (CC BY) license (<http://creativecommons.org/licenses/by/4.0/>).

sciforum-119468: From “no sense of place” to “no sense of the world”: Screen-centered senses and the problem of worldlessness

Zheng Liu

School of History and Culture of Science, Shanghai Jiao Tong University

This article investigates the transformative relationship between screen technologies and human perception through an interdisciplinary framework combining Heidegger’s “revealing/concealing thesis” and postphenomenological analysis. The study systematically examines how screens mediate human-world interactions, reshape sensory engagement, and ultimately reconfigure existential spatial-temporal awareness.

First, Heidegger’s dialectic of technological revealing/concealing provides the foundational lens. While screens reveal hyper-visualized information flows, they simultaneously conceal their material infrastructure and algorithmic operations. This dual process creates an asymmetrical mediation where users engage with curated realities while remaining oblivious to the underlying technical processes that construct them. Postphenomenology further unpacks this mediation by analyzing how screens actively reconfigure perception rather than merely transmitting neutral content.

Second, Albert Borgmann’s “device paradigm” reveals screens’ paradoxical nature as both focal objects and invisible conduits. While screens demand sustained attention through luminous interfaces and infinite scroll dynamics (focal properties), they operate through a “device logic” that obscures their socio-technical complexity. This tension manifests in what we term “distracted immersion”—users become deeply absorbed in screen content while remaining detached from the contextual realities these interfaces mediate.

Third, Meyrowitz’s “no sense of place” thesis is expanded through contemporary screen practices. Location-independent behaviors enabled by smartphones and video conferencing induce spatial dislocation, where physical environments become interchangeable backdrops to screen activities. This de-localization fosters a paradoxical “ubiquitous placelessness”—users operate simultaneously everywhere (digitally) and nowhere (physically), eroding embodied connections to specific locales.

Fourth, postphenomenological distinctions between screen-mediated and screen-centered perception clarify escalating technological influence. Screen-mediated perception enhances human capabilities (e.g., microscopes extending vision), while screen-centered perception creates self-referential ecosystems (e.g., social media feeds prioritizing algorithmic engagement over external reality). The latter generates a closed-loop effect where sensory input and behavioral output both originate from screen interfaces, progressively detaching users from unmediated lifeworld experiences.

Finally, Arendt’s concept of “worldlessness” frames the socio-political consequences. As screens fragment shared reality into personalized digital enclaves, the *sensus communis* essential for public discourse deteriorates. Social media’s filter bubbles and recommendation algorithms exemplify this erosion, replacing communal frames of reference with hyper-individualized perceptual universes. This dissolution of common ground manifests in polarized discourses and the collapse of shared epistemic foundations.

The analysis concludes that screens function as existential mediators reshaping human perception across three dimensions: 1) perceptual (mediated vs. direct experience), 2) spatial (disembodied presence vs. situated embodiment), and 3) social (fragmented collectives vs. communal world-building). These transformations demand critical examination of how screen technologies reconfigure the fundamental structures of human experience and collective existence.



© 2025 by the author(s). Licensee MDPI, Basel, Switzerland. This article is an open access article distributed under the terms and conditions of the Creative Commons Attribution (CC BY) license (<http://creativecommons.org/licenses/by/4.0/>).

sciforum-123355: From GenAI to Human Beings and Back: Assessing Creative Intelligence

Maria Regina Brioschi

University of Milan

This paper aims to provide a thorough understanding of intelligence, based on the pragmatic thought of C.S. Peirce. Generally speaking, pragmatism offers an intriguing perspective on both intelligence and technology. On one hand, it enables us to understand intelligence not in isolation, as if it were disconnected from practice. By starting with the analysis of intelligent behavior and procedures, pragmatism also aids in comprehending human intelligence in continuity with other forms of intelligent behavior that can be found in both living and non-living beings alike. On the other hand, concerning technology the pragmatist approach provides new “conceptual tools” for analyzing and understanding technology beyond the biases and heavy legacies that have characterized philosophy of technology in modern and recent times, such as the polar opposition between an “instrumentalist” view of technology and a “substantivist” one (see Borgoman 1984, Verbeek 2022). This opposition nowadays mirrors the public debates between techno-enthusiasts and technophobes. The former believe that AI, understood as a neutral instrument, may help overcome our limits and empower mankind’s inquiry in various domains. The latter are dominated by worries about the unpredictable and uncontrollable (bad) consequences that the spread of AI will have on our society and species.

The first part of the paper develops such a pragmatist approach and shows its advantages in comparison to current debates, underlying both its philosophical relevance and its interdisciplinary potentiality, in terms of applications. The second part tackles the specific case of GenAI to understand to what extent we should consider it creative and what the differences are with human creative intelligence, as it can be analyzed with reference to both scientific discoveries and art. In this regard, Peirce’s distinction between three kinds of reasoning, namely deduction, induction, and abduction (the latter one being the only creative type, according to the American philosopher), will help determine limits and opportunities in the interactions between GenAI and human intelligence.



© 2025 by the author(s). Licensee MDPI, Basel, Switzerland. This article is an open access article distributed under the terms and conditions of the Creative Commons Attribution (CC BY) license (<http://creativecommons.org/licenses/by/4.0/>).

sciforum-119808: From the Expansion of Economic Rationality to Its Re-Nationalization: Ethics, Purpose, and Geopolitics

Giacomo Maria Arrigo

Università Vita-Salute San Raffaele, Milan, Italy.

Economic rationality has evolved significantly over the past decades, shifting from profit maximization to stakeholder theory and, more recently, to the emergence of *purpose* as a guiding principle for businesses. This shift reflects a broadening of economic intelligence, incorporating ethical, social, and environmental dimensions into corporate decision-making. However, rising geopolitical polarization is reshaping this paradigm: while companies increasingly embrace sustainability and inclusivity, the “cunning of economic reason” (borrowing from Hegel) now appears to align *purpose* with national strategic interests. This raises a fundamental question: is economic rationality truly evolving, or is it being reabsorbed into the logic of power?

This paper adopts a moral philosophy and business ethics perspective, critically engaging with the frameworks proposed by Robert Edward Freeman, Michael E. Porter, and Mark R. Kramer. It examines how the evolution from stakeholder theory to shared value is now undergoing a further transformation, where business purpose is increasingly shaped by geopolitical imperatives. The analysis problematizes these classical models, showing that we are moving beyond them—not by returning to Milton Friedman’s assertion that “the social responsibility of business is to increase its profits,” but by entering a new paradigm where corporate ethics must reconcile global responsibilities with national allegiances. Case studies of SpaceX, Anduril Industries, and Palantir Technologies illustrate how ethical frameworks are being redefined within this tension.

The findings indicate that the boundary between corporations and States is increasingly blurred: major tech firms not only generate economic and social value but also function as geopolitical assets, reinforcing national sovereignty through technological dominance. Palantir, for instance, exemplifies how economic intelligence becomes an extension of State power, merging AI, data control, and defense strategy. In this scenario, economic intelligence is shifting from an emancipatory force to a mechanism of national hegemony. In the uncertainty about whether *purpose* should prioritize the global or the local, lies the core tension between ethics and realism.

This paper highlights a critical dilemma: as economic intelligence becomes increasingly tied to State interests, corporate ethics risk losing their autonomy, turning into instruments of digital *raison d’État*. If shared value is reduced to strategic advantage, can we still consider this an evolution of economic rationality, or is it a strategic involution, where corporate ethics are absorbed into power politics? In an era dominated by AI and big tech, preserving an authentic business ethic requires careful scrutiny of regulatory frameworks and cultural narratives to ensure that *purpose* does not become a new expression of 21st-century realpolitik.



© 2025 by the author(s). Licensee MDPI, Basel, Switzerland. This article is an open access article distributed under the terms and conditions of the Creative Commons Attribution (CC BY) license (<http://creativecommons.org/licenses/by/4.0/>).

sciforum-123414: Gödelian Self-Referential Genomic Information Processing: Complex Self-Other Interactions, Social Cognition and Novelty Production

Sheri Markose

University of Essex

Extended abstract

1. Background

For some 95 odd years since his epochal 1931 paper, Gödel's Incompleteness Theorems (GITs) have been the focus of more heat than light. In this time, there has been very little in the way of evidence that has shown that GITs have any bearing on real-world phenomena.

I aim to unpack recent assertions by Igamberdiev (2021) for why "living systems during evolution continuously realize the proof of Gödel's theorems", or Miller/Torday (2018)'s statement : "As self-referential cognition is demonstrated by all living organisms, life can be equated with the sustenance of cellular homeostasis in the continuous defence of 'self'."

It should be noted that Miller/Torday give centrality to self-referential information processing in genomic systems, specifically to detect and mitigate adversarial changes to self-codes, but they make no reference to the Gödel logic or foundational mathematics of Computation Theory, also known as Recursion Function Theory (RFT). I will draw on my recent work, Markose (2021, 2022), that the Gödel logic and, in particular, the Gödel Sentence, far from being esoteric constructions in the foundations of mathematics, are ubiquitous to genomic information processing in eukaryotes.

Genomic intelligence is a concept introduced in Markose (2021) specifically to characterize the Gödelization of code-based information processing in genomic systems with the distinctive self-referential conditions of Gödel Incompleteness results that appear to have been acquired for complexification over the course of the evolution of multicellular eukaryote life (Markose (2022)). This is a very tall order of computational intelligence that corresponds with so-called Type IV Dynamics in the Wolfram-Chomsky Schema, which are unique in producing novelty. To date, complexity, evolvability, novelty production and 'thinking outside the box' in biology and humans have, for the most, part relied on models of randomness or on statistical white noise error terms. Advances in gene science, neuroscience and security of software systems (cryptography) are beginning to provide evidence that only code-based systems with self-referential recursive machinery in place can produce open-ended novelty production whilst maintaining self-codes hack-free from internal and external malware.

I will first briefly address the two canards that are associated with Gödel Incompleteness Theorems that have posed a stumbling block to the necessary breakthroughs on genomic information processing. The canards are that GITs prove that human cognition is not computational and self-reference leads to paradox. That the latter is not so in the Gödelian setting as unlike the Cretan Liar Paradox *This is False*, Gödel's painstaking application of the Cantor Diagonal Lemma as a means of mechanizing self-reference leads to a theorem. Systems that can mechanize the exit route from listable sets pave the way for syntactic objects (encoded objects) that self-referentially say of themselves a true statement that they cannot be computed if the system is consistent. Influential commentators like Roger Penrose have used GITs to conclude that the human cognition can outstrip what Turing Machines can do. As Rescorla (2020) says : "It may turn out that certain human mental capacities outstrip Turing-computability, but Gödel's incompleteness theorems provide no reason to anticipate that outcome."

2. Unpacking the Evidence for Gödelization of Biology

The digitization of inheritable information in the genome encoded in a near universal alphabet (A,T,C,G/U) has been called the 'algorithmic takeover of biology' by Walker and Davis (2013). The Faustian pact at the genesis of life colourfully portrayed by Freeman Dyson as "the takeover of a replicative digital virus of an analogue metabolism" accords with the perspective of Forterre, Zimmer, Villareal, Koonin and others. This underpins the remarkable fact that, in nature, only life and biology as we know it and the artifacts of genomic intelligence thereof are explicitly code-based software systems.

Significantly, while debunking the idea that the primary source of evolutionary changes arises from random transcription/replication errors, the epochal discovery by Nobel Laureate Barbara McClintock (1984) of viral transposable elements that conduct cut-paste (transposons) and copy-paste (retrotransposon) gives a code based explanation for genomic evolvability, brain plasticity and novel phenotype primarily in eukaryotes. This underscores the truism that primarily only software can change software and also that viral hacking by such internal and external biotic malware is the Achilles heel of genomic digital systems.

While operations on encoded information fall under the purview of Computation Theory and Recursion Function Theory (**RFT**), until recently there was no evidence for how this and Gödel's theorems apply to biology. I will unpack the recent breakthroughs here.

The Gödelization of information processing starts, *firstly*, with unique identifiers or Gödel numbers for digital entities well known in the digital economy and taking the form of bio-peptide unique identifiers including 'zip codes' in organisms as discovered in the Nobel prize-winning work of Blobel (2009). It appears that all signalling in bio-ICT relies on peptide identifiers from transcription factors to neuron-neuron links.

Two other distinctive Gödelian features found in genomic intelligence, using epithets from Hofstadter (1999), are self-reference (**Self-Ref**) or the online machine execution involving the Diagonal operator, and the offline virtual self-representation (**Self-Rep**) of the former. The breakthrough on the significance of these staples of **RFT**, found in textbooks such as Rogers (1967) and Cutland (1980), starts with the insight from Gershenfeld (2012, 2017, Chapter 3, p. 109) as to what the self-referential/Diagonal operator means for biology, where a program g builds machine ϕ to run g (typically denoted as $\phi_g(g)$). Gershenfeld (2012) says that what 21st century digital fabrication is trying to achieve is something biology solved 3.7 billions years ago with the self-assembly programs associated with the ribosome and other transcriptase machinery involved in gene expression for the morphology, somatic identity and regulatory control of the organism. The domain of these self-assembly diagonal machines are the theorems for the eukaryote organism, viz. the self-codes that have to be maintained free from adversarial hacking.

Despite the central role assigned to self-reference for the sentient self in advanced organisms (Gardenfors 2003, Northoff et. al, 2006 , Newen (2018) , Miller et.al., 2018 , etc.), only Tsuda (2014) and Markose (2017, 2021,2022) have noted how the evolutionary development of **Self-Rep** mirror structures as in the Gödel Meta-Representation Theorem (Rogers,1967) is necessary for biotic elements to make statements about themselves, first having self-assembled themselves. Tsuda (2014) makes an explicit observation that unless the mirror Self-Rep recursive structures are in place, it is unlikely that statements about self can be made, let alone make out the other. One can also include here the adversarial other in the form of negation to self-codes. This *offline* embodied **Self-Rep**, which contrasts with no such structures in prokaryotes, was a paradigm shift in the Adaptive Immune System (AIS) some 500 million years ago in the lineage of jawed fish, called the Big Bang of Immunology (Janeway et al. (2005)). This latterly appears as a Mirror Neuron System (MNS) mostly in primate brains first discovered by the Parma Group of neuroscientists.

The **AIS** demonstrates virtual offline mirrored self-representation (**Self-Rep**) in the MHC1 T cell receptors of ~85 % of expressed genes, viz. *halted* machine executions of genomic self-assembly codes that determine the somatic and phenotype identity for the organism. As will be seen, these Self-Repped gene codes in the Thymus, called the *Thymic Self* (Ramon and Faure, 2021) or 'the science of self' (Greenen, 2021), are primarily used to identify the hostile other, viz. negation function operators of non-self antigens. In turn, the Mirror Neuron System (**MNS**) reuses codes of self-actions from the sensory-motor cortex for social cognition and inference regarding conspecifics via virtual simulations in the MNS (Fadiga et al. (1995); Gallese et al. (1996); Rizzolatti et al. (1996)).

It is conjectured that an identical **RFT** machinery is involved in the self–other nexus in both the AIS and MNS. The graphics for **Self-Rep** Mirror Mapping in the AIS and MNS are given in Figure 1 in Markose (2021).

3. Detection of Non-Self Antigen in AIS: New Diversity-Selector Model and Gödel Sentence in Genomic Blockchain Distributed Ledger (BCDL)

The most significant of all breakthroughs here is the one made by the game theorist Binmore (1987) who raised the ‘spectre of Gödel (1931)’ in the form of Gödel’s Liar who will negate or falsify what can be computed/predicted. Binmore effectively mooted the adversarial digital game which is co-extensive with life itself (Markose, 2017, 2021). This constitutes the *fourth* condition of Gödel systems and involves an adversarial agent in the form of a virus or a hacker whose actions cannot be constrained in any way.

The Gödel Incompleteness result that generates the Gödel Sentence permits a code qua biotic element to self-report “I’m under attack”, when it has been hacked/negated by a novel malware. In other words, using the Second Recursion Theorem, the index for the fix point for the novel negation function for the self-code/Theorem for the genomic system is generated. This marks an endogenous exit from listable sets, a necessary condition for the production of novel antibodies, to avoid the irrational state of logical inconsistency of formal systems (Smullyan, 1961). The genomic Gödel Sentence in terms of 21st century nomenclature is a *hash* that helps demonstrate endogenously that the outputs of expressed genes have been maliciously altered.

Indeed, how the **AIS** identifies novel software attacks on own gene codes, g , with malware/parasite negating functions f_p^{-1} , which belongs to an uncountable infinite set that cannot be mechanically listed, is a stupendous case of uber bio-cybersecurity.

The AIS implements an ‘out-of-the-box’ astronomic anticipative search for novel non-self antigens necessary for novel anti-body production and cognition in humans which manifests unbounded proteanism for novel extended phenotypes (Dawkins (1987)) in the form of artifacts outside of ourselves. This facility, found only in the **AIS** relying on the Recombination Activator genes (**RAG** 1 and 2) and in the human brain for neural receptor diversity, runs into orders of magnitude of $10^{20} - 10^{30}$ that exceed the pre-scripted germline genome size many times over.

The Rogers (1967) fixed-point indexes of the Second Recursion Theorem for yet-to-happen f_p^{-1} attacks by the non-self antigens are generated in the **AIS** in the most ingenious fashion: a large number of codes/indexes purported to be of different f_p^{-1} on each **self-repped** g are generated in the T-Cell Receptors. This is the most spectacular case of predictive coding. Suppose that the g_n for the specific tuple $\{f_p^{-1}, g_n\}$ denoted by g^- . When the attack by f_p^{-1} takes place in real time in the periphery involving said pair $\{f_p^{-1}, g_n\}$, the experientially driven peripheral MHC1 receptor must record this and if this ‘syncs’ with the one that was speculatively generated in the thymic MHC1 receptors, two parts of the fixed point come together to construct a genomic Gödel Sentence.

This molecular genomic diversity-selector model follows a unique self-referential blockchain-distributed ledger that is different, in terms of the self-referential design, from man-made **BCDLs** first invented circa 2009. The genomic **BCDL** manifests secure digital and decentralized record-keeping where no internal or external bio-malware can compromise the immutability of life’s building blocks and no novel blocks can be added that are not consistent with extant blocks.

In conclusion, genomic intelligence in vertebrates that has reached its apogee in humans is highly empathic as the conspecific/other is the projection of self; greatly Machiavellian having co-evolved from adversarial viral agents; geared toward unbounded proteanism from the get-go, starting with the transposon-based diversity of **RAG** genes in the immune system and stringently self-regulated by a **BCDL** driven by the principle of autonomy of the life of the organism and an agenda to be hack-free.

Selected References

#Blog Essex University Website: *How we became smart – a journey of discovery through the world of game theory and genomic intelligence* <https://www.essex.ac.uk/blog/posts/2021/10/26/how-we-became-smart>

#Markose, S.M, 2022, Complexification of eukaryote phenotype: Adaptive immuno-cognitive systems as unique Gödelian blockchain distributed ledger, Biosystems, <https://doi.org/10.1016/j.biosystems.2022.104718>

#Markose, S.M, 2021, “Genomic Intelligence as Über Bio-Cybersecurity: The Gödel Sentence in Immuno-Cognitive Systems”, *Entropy*. 23(4), 405; <https://doi.org/10.3390/e23040405>

Markose, S.M, 2017, Complex type 4 structure changing dynamics of digital agents: Nash equilibria or a game with arms race in innovations. *Journal of Dynamics & Games*, doi: 10.3934/jdg.2017015



© 2025 by the author(s). Licensee MDPI, Basel, Switzerland. This article is an open access article distributed under the terms and conditions of the Creative Commons Attribution (CC BY) license (<http://creativecommons.org/licenses/by/4.0/>).

sciforum-121403: How Brooks' behavior-based robots teach us a lesson about the definition of knowledge

Saskia Janina Neumann

Eötvös Lorand University

The world is its own best model (Brooks, 1989)! Or at least it is for Brooks' behavior-based robots. Robots such as Herbert, whose task consisted in stealing empty soda cans from offices, and Squirt, whose task was to follow around noises, prove that complex behavior such as following specific sounds after a specific time interval, avoiding obstacles and real-time recognition (Brooks, 1990) can be achieved without any inner representation of the world. If we believe Hutto and Myin (2013), we human beings are not too different from Herbert or Squirt. Our behavior can be explained without assuming that we mentally represent our world. In my talk, I want to show evidence against this claim. While some of our behavior might be explainable without us representing the world, most of our storing-behavior of the interaction with our world we have is indicative of us representing the world mentally. My claim will be based on empirical evidence stemming from non-linguistically based spreading activation (Barr et al., 2014), false memories (Roediger, 2001), false recognition (Meade et al., 2007) and the processing of visuospatial information (Foster et al., 2017). I will, furthermore, expand this theory and argue for human beings being different in knowledge acquisition to behavior-based robots due to the nature of our memory. While behavior-based robots do not need to represent the world, we human beings have to represent the world in a specific manner to be able to continue to act with the world successfully. The knowledge we acquire is mainly based on our inner representation of the world.

Bibliography

- Barr, R., Walker, J. Gross, J. & Hayne, H. (2014). Age-related changes in spreading activation during infancy, *Child Development* 85(2), 549-563.
- Colombo, M. (2014). Neural representationalism, the Hard Problem of Content and vitiated verdicts. A reply to Hutto & Myin (2013), *Phenom Cogn Sci* 13, 257-274.
- Foster, P. S., Wakefield, C., Prymak, S., Roosa, K. M., Branch, K. K., Drago, V., Harrison, D. W. & Ruff, R. (2017), Spreading activation in nonverbal memory networks, *Brain Informatics* 4, 187-199.
- Hutto, D. D. & Myin, E. (2013). *Radicalizing enactivism, Basic minds without content*, MIT Press.
- Hutto, D. D. (2023). Remembering without a trace? Moving beyond trace minimalism, in: *Current controversies in Philosophy of Memory*, ed: Sant'Anna, A., McCarroll, C. J., Michaelian, K., 59-82.
- Meade, M. L., Watson, J. M., Balota D. A. & Roediger, H.L. (2007). The roles of spreading activation and retrieval mode in producing false recognition in the DRM paradigm, *Journal of Memory and Language* 56, 305-320.
- Roediger, H., Balota, D. A., Watson, J. M. (2001). Spreading activation and the arousal of false memories, *The nature of remembering: Essays in honor of Robert G. Crowder*, (ed) Roediger, H.L., Nairne, J.S., Surprenant, N. A. M., American Psychological Press, 95-115.



© 2025 by the author(s). Licensee MDPI, Basel, Switzerland. This article is an open access article distributed under the terms and conditions of the Creative Commons Attribution (CC BY) license (<http://creativecommons.org/licenses/by/4.0/>).

sciforum-120051: Infoautopoiesis and Intelligence

Jaime F. Cárdenas-García

University of Maryland, Baltimore County

A fundamental issue is whether artifacts can ever gain the status of living beings. In particular, the relevance of meaning-making by machines is an area of study whose relevance cannot be ignored, putting front-and-center the matter of intelligence, whose etymological origin relates to “the ability to understand” [1]. Since the release of ChatGPT on the 30th of November 2022 by OpenAI [2], countless other applications compete in searching through a massive amount of written text by reading millions of articles and books online to produce work that has perfect grammar, correct punctuation, and no spelling mistakes. Multiple uses include writing songs, stories, press releases, guitar tabs, interviews, essays, and technical manuals. In addition, hallucinations are also readily achieved. ChatGPT may be regarded as an autonomous agent, similar to a living organism capable of biological agency, which acquires, uses, processes, communicates, and acts on information [3]. This brings into focus the use of information as the glue that makes possible the examination of intelligence in living beings and their artifacts. The purpose of this presentation is to critically examine information in this role.

The long history of information uncovers an elusive concept that needs clarification [4, 5], and involves a dichotomy that needs resolution. For some, information is considered an absolute quantity of the Universe in addition to matter and/or energy, whose existence is predicated upon a postulate which some consider sufficient to bring into existence [6-11]. For others, it is a relative quantity/quality, ‘a difference which makes a difference’ [12], where “The essence of this definition is that information is something which is generated by a subject. Information is always information for “someone”; it is not something that is just hanging around “out there” in the world” [13]. This implies that there is no information outside living beings interacting with their environments. Clearly, the more reliable choice to use to perform a more detailed assessment of information is the one not dependent on the enunciation of a postulate, but rather, on firsthand observation.

The result of performing this assessment leads to infoautopoiesis, the self-referential, recursive, and interactive process of self-production of information [14]. Not only is it possible to create a model of a general organism-in-its-environment, but it is also possible to also define the roles of semantic (endogenous) and Shannon/syntactic (exogenous) information [15]. Semantic information creation is motivated by the individuated satisfaction of physiological and relational needs in order to make the external environment meaningful. Semantic information is inaccessible except through external Shannon/syntactic information expressions using language, gestures, pictographs, musical instruments, sculptures, writing, coding, etc. Shannon/syntactic information is a metaphor for the creation of all our artificial artifacts in the arts and sciences which surround us [16-18]. In summary, infoautopoiesis enables an explanation of how meaning-making comes about, and its use as a general framework to answer the key question: what is the connection between human and machine intelligence?

References

1. Da Silveira, T.B.N. and H.S. Lopes, *Intelligence across humans and machines: a joint perspective*. Frontiers in Psychology, 2023. **14**.
2. Roose, K. *The Brilliance and Weirdness of ChatGPT*. New York Times, 2022. DOI: <https://www.nytimes.com/2022/12/05/technology/chatgpt-ai-twitter.html>.
3. Dodig-Crnkovic, G. and M. Burgin, *A Systematic Approach to Autonomous Agents*. Philosophies, 2024. **9**(2): p. 44.
4. Capurro, R. and B. Hjørland, *The concept of information*. Annual Review of Information Science and Technology, 2003. **37**(1): p. 343-411.
5. Capurro, R., Past, present, and future of the concept of information. TripleC, 2009. **7**(2): p. 125-141.

6. Wheeler, J.A. Sakharov revisited: "It from Bit". in Proceedings of the First International A D Sakharov Memorial Conference on Physics, May 27-31. 1991. Moscow, USSR: Nova Science Publishers, Commack, NY.
7. Stonier, T., *Information and Meaning - An Evolutionary Perspective*. 1997, Berlin Heidelberg New York: Springer-Verlag.
8. Yockey, H.P., *Information theory, evolution, and the origin of life*. 2005, Cambridge, UK: Cambridge University press.
9. Lloyd, S., *Programming the Universe*. 2006, New York, NY: Alfred A. Knopf.
10. Floridi, L., *Information: A Very Short Introduction*. 2010: Oxford University Press.
11. Vedral, V., *Decoding Reality - The Universe as Quantum Information*. 2010, Oxford, UK: Oxford University Press.
12. Bateson, G., *Steps to an ecology of mind; collected essays in anthropology, psychiatry, evolution, and epistemology*. Chandler publications for health sciences. 1978, New York: Ballantine Books. xxviii, 545.
13. Hoffmeyer, J., *Signs of meaning in the universe*. 1996: Bloomington : Indiana University Press, [1996] ©1996.
14. Cárdenas-García, J.F., The Process of Info-Autopoiesis - the Source of all Information. *Biosemiotics*, 2020. **13**(2): p. 199-221.
15. Burgin, M. and J.F. Cárdenas-García, A Dialogue Concerning the Essence and Role of Information in the World System. *Information*, 2020. **11**(9): p. 406.
16. Cárdenas-García, J.F., Info-Autopoiesis and the Limits of Artificial General Intelligence. *Computers*, 2023. **12**(5): p. 102.
17. Cárdenas-García, J.F., *Information is Primary and Central to Meaning-Making*. *Sign Systems Studies - Special Issue on Contemporary Applications of Umwelt Theory*, 2024. **52**(3/4): p. 371-393.
18. Cárdenas-García, J.F., *Syntactic Touch: A Probe*. *New Explorations: Studies in Culture and Communication*, 2024. **4**(2).



© 2025 by the author(s). Licensee MDPI, Basel, Switzerland. This article is an open access article distributed under the terms and conditions of the Creative Commons Attribution (CC BY) license (<http://creativecommons.org/licenses/by/4.0/>).

sciforum-124989: Intelligence and Consciousness in Natural and Artificial Systems

David Gamez

Department of Computer Science, Middlesex University, London, UK

Considerable progress has been made with the development of systems that can drive cars, play games, predict protein folding and generate natural language. These systems are described as intelligent and there has been a great deal of talk about the rapid increase in artificial intelligence and its potential dangers. However, our theoretical understanding of intelligence and ability to measure it lag far behind our capacity for building systems that mimic intelligent human behaviour. There is no commonly agreed definition of the intelligence that AI systems are said to possess, nor has anyone developed a practical measure that would enable us to compare the intelligence of humans, animals and AIs on a single scale. This talk addresses these problems by clarifying the nature of intelligence and outlining a new algorithm for measuring intelligence that can be applied to any system.

The first part of the talk starts with a discussion of previous definitions of intelligence. It then argues for a close link between prediction and intelligence and addresses two misconceptions about intelligence. The first is that people often think that humans have a general form of intelligence that has the same level in all environments. This belief motivates the idea that we could develop machines with artificial general intelligence (AGI). However, human intelligence often fails when it is confronted with environments that are significantly different from the natural world, such as high-dimensional numerical spaces. A second issue is that people naively assume that they directly apply their intelligence to the physical world. However, we can only be intelligent about things that are revealed to us through our senses, and people, animals and artificial systems have very different sensory experiences. So, it is much more accurate to say that agents apply their intelligence to their perceived environment, or *umwelt*.

The second part of the talk explores the measurement of intelligence. Previous work in this area includes IQ, *g* and universal measures, such as compression tests and algorithms based on goals and rewards. To address the limitations of previous measures, I have developed a new algorithm for measuring predictive intelligence that is based on an agent's internal state transitions. Experiments have been done to test this algorithm, and it has many potential applications in AI safety and the comparative study of intelligence.

The talk concludes with some reflections on the relationships between intelligence and consciousness. It is commonly assumed that there is a close relationship between intelligence and consciousness in biological systems. However, this correlation might not exist in artificial systems, who could be highly intelligent with low levels of consciousness, or highly conscious with low levels of intelligence. In the future we might be able to use algorithmic measures of consciousness, such as information integration theory (IIT), and universal measures of intelligence to systematically study the relationships between intelligence and consciousness in natural and artificial systems.



© 2025 by the author(s). Licensee MDPI, Basel, Switzerland. This article is an open access article distributed under the terms and conditions of the Creative Commons Attribution (CC BY) license (<http://creativecommons.org/licenses/by/4.0/>).

sciforum-122856: Intelligence and The “Hard Problem” of Consciousness—a short analysis

Marcus Abundis

Formerly: Organizational Behavior (GFTP), Graduate School of Business, Stanford University (March 2011)

SUMMARY: This paper situates the study of intelligence in relation to David Chalmers’s anti-scientific Hard Problem, suggesting that evolution by means of natural selection (EvNS) and information theory are likely useful scientific approaches.

ABSTRACT: As a prelude to ‘intelligence studies’, one must cover the long-held claim of an imagined Hard Problem—that a grasp of the human mind lies wholly beyond scientific views.

The formal study of human mentation surely remains challenging, as persistent ‘material–immaterial’ analogues are represented in the literature as a “symbol grounding problem”, “solving intelligence”, a missing “theory of meaning”, and more. The above Hard Problem claim thus has high intuitive appeal, lying open since the time of pre-socratic philosopher Anaxagoras. Hence, issues of ‘subjective phenomena’ –in relation to informatic intelligence studies—still hold sway in many corners, being unresolved.

But to say that a Hard Problem eternally surpasses science lacks equal intuitive appeal. Moreover, the literature shows that the close study of the Hard Problem is rare. Instead, most researchers continue with an ‘intuitive vein’, often claiming that 1) the Hard Problem is a plainly absurd view unworthy of study, or 2) it is an innately intractable issue beyond practical study, where neither side ever offers much actual clarifying detail.

As such, this paper takes a different approach—that of firmly assessing the Hard Problem’s original statement(s) contra one specific scientific role: evolution by means of natural selection (EvNS). The paper closely examines the logic behind the Hard Problem’s claims, as seen in the literature over the years.

The paper’s analysis ultimately shows that the Hard Problem’s logic is deeply flawed, especially in relation to EvNS—with the implication that EvNS still remains available for explaining/exploring consciousness. Moreover, this study suggests that an ‘information theory’ approach is likely best for addressing an extant material–immaterial divide—paper available.

KEYWORDS: science; philosophy; information; information theory; information science; hard problem; intelligence; consciousness; natural selection; zombies



© 2025 by the author(s). Licensee MDPI, Basel, Switzerland. This article is an open access article distributed under the terms and conditions of the Creative Commons Attribution (CC BY) license (<http://creativecommons.org/licenses/by/4.0/>).

sciforum-123289: Intelligence as a second-order virtue: changing attitudes for successful interactions in digital environments.

Francisco Miguel Macías-Pozo

University of Málaga

From an instrumental perspective, intelligence is understood as a problem-solving ability or error avoidance (Holm-Hadulla et al. 2022). It is used to fulfil goals through the design and selection of good strategies or attitudes. The success can be understood in terms of the interests and expectations of each person. But when referring to inquiry, authors like Kruglanski and Boyantki (in Matheson and Vitz, 2014), and Tanesini (2021) among others, two main goals are generally shared: social recognition –direct or indirect through other kinds of non-epistemic achievements– or truth-conduciveness –including avoiding falsehoods and not only finding truths.

Intelligence makes use of our perspective (as a set of beliefs) and epistemic evaluations for choosing strategies, believing those intentions will match their objectives at least better than others. So true beliefs are part of intelligence, as they are doxastic attitudes not only considered approximately true (or supposed as true) description of states of things but are shaping our perspective and evaluation of objects and propositions, like hypotheses and, with differences, like emotions.

I will defend that intelligence is a second-order epistemic virtue because it allows epistemic subjects to choose correctly how to display their character traits to produce good effects, generally, satisfying their own interests. In this, I follow a normative contextualism approach for traits, considering them virtuous or vicious depending on several situational factors, such as the kind of motives or effects around the actions guided by the attitudes and decisions (Kidd et al. 2021, 82). Using conspiracist echo chambers as a case study, this talk will delve into how being inside an echo chamber creates a scenario where intelligence plays a fundamental role, as in our previous beliefs, in shaping strategies and attitudes for dealing with the problems of this epistemic and emotional environment.

Thus, this work offers support for character-based and decision-making theories to be a constitutive part of studies about intelligence. For that, it is essential to provide empirical evidence about the circumstances influencing the psychological processes and their outcomes while being immersed in social media environments, which is, according to Levy and Mandelbaum (in Matheson and Vitz, 2014), a new set of different from our past ones, for which humans are not well-equipped.

The research question is: how can people choose smartly what to do when they find themselves inside an echo chamber? My thesis will be that intelligence is a second-order virtue, in the sense that every person has to display their traits as virtues while countering her traits to avoid the production of bad effects, becoming vices. Intelligence's main functions, if so, consist of analysing the relationships between internal and external situational factors, and coordinating actions for achieving goals in each situation.

Key words: virtue epistemology; character traits; intelligence; decision theory; echo chambers

Philosophical issues addressed in the work: the debates between situationism and virtue epistemology, and inside the latter, between responsibilism and reliabilism. Also, I address the issues related to echo chambers, as the recognition of being inside one, how to deal with the situation and which could be the proper actions, depending on internal and external factors.



© 2025 by the author(s). Licensee MDPI, Basel, Switzerland. This article is an open access article distributed under the terms and conditions of the Creative Commons Attribution (CC BY) license (<http://creativecommons.org/licenses/by/4.0/>).

sciforum-123384: Intelligence as Typological Cognition: Revisiting Jungian Functions for Human and Artificial Minds

Zijian Ding

School of Philosophy, Psychology and Language Sciences, University of Edinburgh

Conventional definitions of intelligence have been shaped by a rationalist tradition, where intelligence is reduced to logical reasoning, information processing, and performance optimization. In this view, intelligence is often equated with a monolithic thinking function. However, this understanding overlooks both the diversity and the inner tensions within human cognition. Recent psychological and neurocognitive research supports a more complex view: Daniel Goleman's Emotional Intelligence theory highlights emotional regulation as integral to intelligence; Daniel Kahneman's dual-process model shows that intuitive, fast thinking is pervasive in reasoning; and Dario Nardi's neurophysiological studies identify differentiated brain activity patterns corresponding to Jungian functions.

Carl Jung's theory of psychological types provides an early and profound contribution to this line of research. Jung proposed that cognition is structured not only by four basic functions—Thinking, Feeling, Intuition, and Sensation—but also by two opposing orientations: introversion (inner-directed) and extraversion (outer-directed). Crucially, Jung argued that thinking itself can be divided into Introverted Thinking (Ti), which is concerned with internal coherence and system-building, and Extraverted Thinking (Te), which is focused on external data and pragmatic results. Moreover, Feeling is conceptualized as a rational function that evaluates situations based on subjective values (Fi) or social harmony (Fe). The tension between introverted and extraverted orientations means that developing one function tends to suppress its opposite. For instance, a Ti-dominant person may sacrifice practical results for internal conceptual clarity, while a Te-dominant agent may disregard subjective depth for external efficiency. This typological framework reveals that intelligence is never uniform—it is informed by differentiated and sometimes conflicting cognitive functions, each shaping distinct ways of engaging with reality or constructing inner cognitive systems. Intelligence, from this perspective, is never uniform or neutral—it is always selective, structured, and shaped by competing cognitive priorities.

Jung's eight-function model offers a compelling map of intelligence as a multi-dimensional phenomenon, with each function corresponding to a distinct mode of information processing and world engagement: four rational/judging functions: logical-analytic (Ti, introverted thinking), pragmatic-executive (Te, extraverted thinking), subjective value-based (Fi, introverted feeling), social-relational (Fe, extraverted feeling), and four irrational/ perceiving functions: intuitive-visionary (Ni, introverted intuition), creative-divergent (Ne, extraverted intuition), experiential memory-based (Si, introverted sensing), and perceptual action-oriented (Se, extraverted sensing).

These functions shape not only individual personality but also collective human knowledge production, including philosophy, science, and technology. No theory or worldview is purely objective; all cognition operates through functional preferences and typological biases. This insight is especially relevant for Artificial Intelligence (AI), as AI systems inherit human cognitive biases through data input, models, and design logic.

From this perspective, intelligence is not merely a computational capacity but a multi-faceted construct. Current AI development shows clear functional imbalances: e.g., while excelling in Te-like data optimization and Se-like sensory processing, AI systems remain underdeveloped in Fi-style ethical judgment and Si-style experiential learning. For both human cognitive research and AI development, typological cognition offers a necessary tool for better understanding cognitive diversity, developing all-round intelligence, and guiding future AI towards more human-like, ethical, and context-sensitive architectures.



© 2025 by the author(s). Licensee MDPI, Basel, Switzerland. This article is an open access article distributed under the terms and conditions of the Creative Commons Attribution (CC BY) license (<http://creativecommons.org/licenses/by/4.0/>).

sciforum-120238: Intelligent behaviour as adaptive control guided by accurate prediction

Nina Poth^{1*}, **Trond A. Tjøstheim**² and **Andreas Stephens**³

¹ Department of Philosophy of Mind and Language at the Faculty of Philosophy, Theology and Religious Studies at Radboud University Nijmegen

² Department of Philosophy, Cognitive Science, Lund University, Sweden

³ Department of Philosophy, Theoretical Philosophy, Lund University, Sweden

What is intelligence and why is it needed? Current research in the cognitive and life sciences presents only fragmented views on this question. We examine the recent proposal that intelligence can be understood as accurate prediction, and develop it behaviourally to include adaptive control. We argue that this framework, albeit already applied to diverse domains of intelligence research, needs conceptual clarification. We specifically refine the notions of “accuracy” in terms of practically exploitable representation and “prediction” in terms of action-oriented probabilistic inferences at the subpersonal level of information processing. We then argue that this interpretation does not fully explain the mechanisms of intelligence in artificial systems, but that it nevertheless provides a useful analysis. Firstly, this view allows us to demarcate cognition from intelligence, the latter of which comes out as a more sophisticated form of efficient and robust information processing that may be realised in systems that are non-cognitive (e.g., bacteria and large-language models). Secondly, this view provides a unified platform for researchers from different disciplines to integrate their diverse theoretical perspectives, share fragmented data, and effectively discover intelligence in non-human systems. We also explore some obstacles and avenues to scaling the view up to adequately capture intelligence in social and collective systems.

Our framework ties the different notions of intelligence together in a natural way by integrating them into a control theory framework of adaptive control and model-based reinforcement learning. We posit that memory, knowledge, experience, and understanding can be interpreted as having to do with an adaptive controller’s predictive model, where memory pertains to the ability to store sensory traces, knowledge indicates that memory traces are structured and consolidated to afford effective prediction, and experience indicates the sensory traces in the context of an agent being situated and continuously sampling the world, learning both its environment’s structure and statistics, as well as the statistics of its own body. So long as it is the case that there is uncertainty both about the state of the environment and which behaviour best achieves the agent’s goals, the agent needs to make decisions and choices based on a value system where some outcomes are better or worse than others. Together, these constitute the ability of the agent to make judgements, weighting and prioritising behaviours and outcomes in relation to its internal needs.



© 2025 by the author(s). Licensee MDPI, Basel, Switzerland. This article is an open access article distributed under the terms and conditions of the Creative Commons Attribution (CC BY) license (<http://creativecommons.org/licenses/by/4.0/>).

sciforum-120318: Jean Piaget and Objectivity—Genetic Epistemology's Place in a View from Nowhere

Mark A Winstanley

The most natural expression of the aim of objective knowledge according to Thomas Nagel is 'we must get outside of ourselves, and view the world from nowhere within it' [1]. Taken literally, however, he finds it unintelligible. Realistically, we cannot in fact get outside of ourselves; we can only hope to achieve a more detached conception by relying 'less and less on certain individual aspects, and more and more on something else, less individual, which is also part of us' [1]. This 'self-transcendent conception should ideally explain (1) what the world is like; (2) what we are like; (3) why the world appears to beings like us in certain respects as it is and in certain respects as it isn't; (4) how beings like us can arrive at such a conception.' [1]

Nagel's self-transcendent conception of objective knowledge is dynamic, involving a dialectic between changes in knowledge of (1) and (2) to form an explanation of (3). In modern times, epistemological problems often occur due to advancements in the human sciences; however, Nagel laments that (4) is often lacking: 'We tend to use our rational capacities to construct theories, without at the same time constructing epistemological accounts of how those capacities work. Nevertheless, this is an important part of objectivity. What we want is to reach a position as independent as possible of who we are and where we started, but a position that can also explain how we got there.' [1, e.g., 2]

Jean Piaget's research, especially his research on the psychogenesis of concepts, contributed to (2); however, this was done to epistemological ends. He conceived of a scientific epistemology founded solely on development, and the psycho- and historiogenesis of knowledge were its methodological pillars. It is known as 'genetic epistemology' but represented a research programme rather than a finished theory [3]. Nevertheless, after decades of research, Piaget concluded that objective knowledge is a process rather than a state [4]. More importantly for the present purposes, however, genetic epistemology also provided an explanation for (4).

In this paper, I argue that Piaget explains how beings like us develop rational capacities to construct theories. Having first given my reasons for choosing Nagel's conception of objectivity in the Introduction, I proceed to characterise 'theory' and show that the concept has both logical and algebraic descriptions under Logical and Algebraic Descriptions of Theories. Under Development of Our Rational Capacities, I then briefly sketch Piaget's account of how our rational capacities develop. This development culminates in hypothetico-deductive thought, and, under The Structure of Hypothetico-Deductive Thought, I set out the interpropositional grouping, which characterises our rational capacities at this stage of development. Under The Interpropositional Grouping: A Canonical Theory, I compare the interpropositional grouping with the algebraic description of theories and conclude that the interpropositional grouping actually represents an archetypical theory. Finally, I conclude by briefly summarising my findings before locating genetic epistemology in a view from nowhere.

References

1. Nagel, Thomas. 1986. *The View From Nowhere*. New York; Oxford: Oxford University Press, USA.
2. Johnson-Laird, Philip N. 2008. *How we Reason* - Oxford Scholarship. October 23.
3. Piaget, Jean. 1950. *Introduction à l'épistémologie génétique*. (I) *La pensée mathématique*. Electronic version from Fondation Jean Piaget pour recherches psychologiques et épistémologiques. Pagination according to 1st edition 1950. Vol. 1. 3 vols. Paris: Presses Universitaires de France.
4. Piaget, Jean, and Rolando Garcia. 1989. *Psychogenesis and the History of Science*. Translated by Helga Feider. New York: Columbia University Press.



© 2025 by the author(s). Licensee MDPI, Basel, Switzerland. This article is an open access article distributed under the terms and conditions of the Creative Commons Attribution (CC BY) license (<http://creativecommons.org/licenses/by/4.0/>).

sciforum-121502: Key Methodologies in Intelligence Studies: Techniques, Skills, and Integration

Dr. Sunila Atul Patil * and Amruta Nilesh Patil

Department of Pharmaceutical Chemistry, P.S.G.V.P.M's College of Pharmacy, Shahada 425409, Maharashtra, India

Intelligence studies utilize a variety of methodologies to gather and analyze information, each requiring specialized skills and techniques. These methods include Open Source Intelligence (OSINT), which involves collecting data from publicly accessible sources such as websites and social media; Human Intelligence (HUMINT), which focuses on obtaining information through direct interaction with human sources; Signals Intelligence (SIGINT), which involves intercepting electronic communications; Geospatial Intelligence (GEOINT), which uses satellite imagery and geospatial data; and Cyber Intelligence (CYBINT), which focuses on monitoring cyber activities and threats. Each methodology employs distinct techniques, such as data mining, signal interception, and spatial analysis, and demands proficiency in skills like data processing, interpersonal communication, and technical analysis. These approaches are often integrated to provide comprehensive intelligence for decision-making and strategic planning.



© 2025 by the author(s). Licensee MDPI, Basel, Switzerland. This article is an open access article distributed under the terms and conditions of the Creative Commons Attribution (CC BY) license (<http://creativecommons.org/licenses/by/4.0/>).

sciforum-119830: Knowledge as an emergent effect of the complexity of large language models

Rafal Maciag

Jagiellonian University, Institute of Information Studies

The following abstract contains a proposal for interpreting the phenomenon of knowledge that appears during the exploitation of large language models (LLMs). As the ontological context of this knowledge, language is proposed, which is also a social phenomenon. Knowledge is interpreted here as an emergent effect of the complex system, which is the working LLM. To explain this effect, we propose to build a trajectory for each input token which is based on the theory of discursive space. These trajectories traverse the manifold that contains all the computational spaces that each token traverses. This reasoning is justified by the nonlinear and nondeterministic nature of LLMs. This task requires at least a preliminary understanding of the descriptive category of what knowledge is. Evidence of the importance of knowledge in the context of LLMs is also provided by research, e.g., Fierro et al., Heersmink et al., Peterson, Kim, and Thorne.

The approach to knowledge here is pragmatic, which means that it is assumed that the form of retention/articulation of knowledge is broadly understood linguistic utterances. Language is a product of social circumstances and therefore remains permanently determined historically and territorially. It also produces deep, abstract phenomena capable of transferring knowledge, which primarily include discourse. This approach was proposed by Michel Foucault.

The result of LLM training is a stable computational structure of great complexity, which is illustrated by the large number of parameters. The resulting structure, which is a composition in time of many multidimensional computational spaces, is not strictly deterministic.

In the process of processing semantic inferences in a model based on a stable, trained set of parameters, a specific "phase transition" occurs from the numerical level to the epistemic level (of knowledge as a phenomenon). The concept of emergence can be used to describe this "transition", and the hypothesis can be formulated that knowledge is an emergent effect of a complex LLM system, in the sense given to emergence by Philip Clayton, developing the definition by el-Hani and Pereira.

The emergence problem can theoretically be solved by a model of the trajectory of semantic signatures, i.e., vectors in different spaces that build the LLM, according to the order of calculations.

The formal way of analyzing knowledge in LLMs is to reconstruct the trajectories of semantic signatures in a manifold, which would consist of all the computational spaces of the model. These trajectories would be a representation of the particular model's knowledge and would be visible to the outside as an emergent phenomenon of knowledge. This approach would constitute an extended application of the definition of knowledge proposed in the theory of discursive knowledge by Rafal Maciag: "knowledge is a set of states of gnosemes in an n-dimensional manifold that can be interpreted locally as a knowledge space".

Philosophical issues:

1. Epistemological issues—the problem of knowledge as a social phenomenon;
2. Ontological issues—the problem of complex systems and accompanying phenomena, e.g. emergence;
3. Hermeneutical issues—since the subject of the research is an artificial text, the hermeneutical context becomes important.



© 2025 by the author(s). Licensee MDPI, Basel, Switzerland. This article is an open access article distributed under the terms and conditions of the Creative Commons Attribution (CC BY) license (<http://creativecommons.org/licenses/by/4.0/>).

sciforum-123114: Rethinking intelligence: the problems of the representational view in the era of LLMs

Johan Largo

Institute of Philosophy, University of Luxembourg

Large language models (LLMs), such as ChatGPT, seem to exhibit human-level intelligence on a wide range of cognitive tasks. This raises the following questions: is such behavior enough to determine the kind of intelligence current AI systems possess? Is this seemingly intelligent behavior enough to attribute to AI the same kind of intelligence we commonly attribute to humans? Pioneer studies (Bubeck et al., 2023) suggest that systems such as ChatGPT start to come close to the generality observed in human intelligence. Such approach has been characterized as using a methodology of “black-box” interpretability. However, “black-box” methodologies are commonly criticized because they ignore the inner functioning of a model, by just analyzing the inputs and outputs of a given system, and because they assume, arguably without good reasons, that the same kinds of tasks that measure human intelligence might also shed some light into LLMs’ intelligence. The argument, inspired by Dretske (1993) and recently used by Grzankowski (2024), about how to find traces of “real intelligence” in such systems and why “black-box” approaches are potentially misleading has the following form:

1) Intelligence depends on a) mental representations of the correct kind (semantic kind) and b) those same relevant mental representations have to be “used” by the system in order to cause behavior.

2) Black-box approaches to LLMs fail to account for a) and b).

3) Then, we have no reasons to think that there is intelligence in LLMs (at least according to black-box approaches).

The main goal of my paper is to argue against premise 1) by criticizing both a) and b) as conditions for intelligence attribution. I remain slightly neutral about 2) and 3).

As a first step, I problematize the concept of mental representation itself as a scientific or natural kind. Consequently, I also show that there are fruitful candidates for definitions of intelligence that do not appeal to mental representations at all. More specifically, in the same spirit as arguments proposed by Ramsey (2017), I show how there can be a useful “divorce” between cognitive science and mental representation, specifically when it comes to intelligence. This objection also draws on recent objections to mental representations as scientific kinds (Facchin, 2023) or as natural kinds (Allen, 2017). I show that, in the relevant sciences, different definitions of intelligence have been proposed which are compatible with a more deflationary approach to mental representation. This small survey includes popular characterizations in psychology (Sternberg, 2019; Gardner, 2011), neuroscience (Haier, 2023; Duncan, 2020) and AI itself (Russell, 2016; Brooks, 1995). Provisionally, the argument in this first section has the following form:

1) Intelligence requires mental representations.

2) Mental representation is a scientifically problematic concept.

3) Then, intelligence is a scientifically problematic concept. (But actual scientific views do not seem to require mental representations in any case.)

As a second step, I use the distinction between vehicle and content to problematize the causal efficacy of mental content, and argue, borrowing from Egan (2020), that contents and their causal/explanatory role are pragmatically but not metaphysically constrained. Additionally, I argue that, even assuming that representations in LLMs have some structural resemblance to their targets, which makes them good candidates of non-arbitrary content determination, they still suffer of potential causal inefficacy (Nirshberg, 2023) and, therefore, should be considered as mere epiphenomena (Baysan, 2021). Following this line of thought, I argue that even scientific methods like probing LLM representations rely on pragmatic attributions, as it has been argued in neuroscience regarding the use of probes for decoding mental representations (Cao, 2022). Such a caveat might allow us to think of mental representation as an important explanatory concept in

AI's behavior, without engaging with the metaphysical problems that the concept commonly implies, particularly regarding its causal efficacy. Consequently, a useful concept of intelligence should drop the causal requirement of contents without fully trivializing the role of representations in behavioral explanations. This second argument has roughly the following structure:

- 1) Intelligence requires mental contents to be causal explanations of behavior.
- 2) Mental contents are not causal explanations of behavior.
- 3) Then, intelligence is an incoherent concept. (Not even applicable to humans.)

Finally, I present some speculative considerations about a more deflationary concept of intelligence that may have the potential advantage of being scientifically productive and avoid severe metaphysical problems. The general conclusion of my paper is two-sided: it has a negative and a positive conclusion. On the negative side, I claim that, at least until we have a more robust account of mental representation in science and philosophy, we should think of intelligence without necessarily requiring mental representation or its causal powers. On a more positive side, we can evaluate intelligence in terms of more operational criteria, for example, behavioral success, mechanistic complexity and learning conditions, allowing us to take seriously the intelligence that we observe in generative systems. Such a conclusion does not imply that generative systems are categorically intelligent or even more intelligent than humans just because they behave as such, at least with respect to certain cognitive tasks. The mere fact that their mechanistic complexity is limited and not as efficient as the one of the human brain, along with the fact that such systems require vastly more instances and are less flexible in learning when compared to human children, are sufficient reasons to think of them in very relevant respects as less intelligent than humans. Nevertheless, LLMs invite rethinking a concept that, despite being vague, may be more scientifically productive by excluding unjustified anthropocentric requirements, and may have better scientific grounds by being operationally applicable to a wide range of systems, from simple and biological to complex and artificial ones.

References

- Allen, C. (2017). On (not) defining cognition. *Synthese*, 194(11), 4233–4249. <https://doi.org/10.1007/s11229-017-1454-4>
- Baysan, U. (2021). Rejecting epiphobia. *Synthese*, 199(1–2), 2773–2791. <https://doi.org/10.1007/s11229-020-02911-w>
- Brooks, R. A. (1995). Intelligence without reason. In L. Steels & R. A. Brooks (Eds.), *The artificial life route to artificial intelligence: Building embodied, situated agents* (pp. 25–81). Lawrence Erlbaum Associates, Inc.
- Bubeck, S., Chandrasekaran, V., Eldan, R., Gehrke, J., Horvitz, E., Kamar, E., Lee, P., Lee, Y. T., Li, Y., Lundberg, S., Nori, H., Palangi, H., Ribeiro, M. T., & Zhang, Y. (2023). Sparks of Artificial General Intelligence: Early experiments with GPT-4. <http://arxiv.org/abs/2303.12712>
- Cao, R. (2022). Putting representations to use. *Synthese*, 200(2). <https://doi.org/10.1007/s11229-022-03522-3>
- Dretske, F. (1993). Can Intelligence Be Artificial? In *An International Journal for Philosophy in the Analytic Tradition* (Vol. 71, Issue 2).
- Duncan, J., Assem, M., & Shashidhara, S. (2020). Integrated Intelligence from Distributed Brain Activity. In *Trends in Cognitive Sciences* (Vol. 24, Issue 10, pp. 838–852). Elsevier Ltd. <https://doi.org/10.1016/j.tics.2020.06.012>
- Egan, Frances (2020). A Deflationary Account of Mental Representation. In Joulia Smortchkova, Krzysztof Dołęga & Tobias Schlicht (eds.), *What Are Mental Representations?* New York, NY, United States of America: Oxford University Press.
- Facchin, M. (2023). Why can't we say what cognition is (at least for the time being). *Philosophy and the Mind Sciences*, 4. <https://doi.org/10.33735/phimisci.2023.9664>
- Gardner, H. (2011). *Frames of mind: the theory of multiple intelligences*. Basic Books.
- Grzankowski, A. (2024). Real sparks of artificial intelligence and the importance of inner interpretability. *Inquiry*, 1–27. <https://doi.org/10.1080/0020174X.2023.2296468>
- Haier, R. J. (2023). *The neuroscience of intelligence*. Cambridge University Press.
- Li, K., Hopkins, A. K., Bau, D., Viégas, F., Pfister, H., & Wattenberg, M. (2022). Emergent world representations: Exploring a sequence model trained on a synthetic task. *arXiv preprint arXiv:2210.13382*.

Nirshberg, G. (2023). Structural Resemblance and the Causal Role of Content. *Erkenntnis*. <https://doi.org/10.1007/s10670-023-00699-y>

Ramsey, W. (2017). Must cognition be representational? *Synthese*, 194(11), 4197–4214. <https://doi.org/10.1007/s11229-014-0644-6>

Russell, S. (2016). Rationality and intelligence: A brief update. *Fundamental issues of artificial intelligence*, 7–28.

Sternberg, R. J. (2019). The Augmented Theory of Successful Intelligence. In *The Cambridge Handbook of Intelligence*. Cambridge University Press. <https://doi.org/10.1017/9781108770422>



© 2025 by the author(s). Licensee MDPI, Basel, Switzerland. This article is an open access article distributed under the terms and conditions of the Creative Commons Attribution (CC BY) license (<http://creativecommons.org/licenses/by/4.0/>).

sciforum-119457: The Cognition System Theory

Jevgenija Sivoronova ^{1,*} and Aleksejs Vorobjovs ²

¹ PhD Candidate, Education and Psychology Department, Faculty of Humanities and Social Sciences, Daugavpils University, Daugavpils, Latvia

² Professor Emeritus, Education and Psychology Department, Faculty of Humanities and Social Sciences, Daugavpils University, Daugavpils, Latvia

We state that when we think about cognition, we must have an ideal image or images of it in our minds at every moment. This paper proposes a cognition model as the ideal image for discussion. It is more of a theoretical or hypothetical perfect model than an entirely abstract one, as it can be completely concrete. This is so because it represents a theoretical position we hold almost meditatively, serving as a mental foundation critical for thinking about any matter. It is a perspective taken by us from an observer's position, which helps us continually reflect on cognition, ourselves, objects of thinking, and knowledge. Ultimately, cognition from the observer's position is thinking about cognition. Thus, we propose one possible approach that takes a holistic view, which we conceive as the theory of the cognition system. Taking this approach, we examine the co-construction of ontologies and epistemologies and employ a three-tier methodology to create a position from which a thing, concept, or matter can be grasped and thought about.

This system represents a complex and dynamic structure of ideas and models the concept of cognition by defining three interconnected levels: philosophical, general scientific, and specific scientific modelling. Based on these premises, we create the position by first determining the philosophical grounds, which involve reflecting on core ideas, shaping a paradigm, and formulating principles and approaches. Next, we define the general modelling principle, outlining the forms of knowledge and methods used to model cognition and explore its elements, whether as a phenomenon, process, or tangible object, and us as the empirical subjects and consciousness. Finally, we select from various scientific approaches, concepts, theories, and research principles to observe cognition through specific entities. This three-level conceptualisation is reflected in the presentation's three sections. Then, the fourth section presents the synthetic culmination, answering the question about the observer's nature and its conceivable actions and potentials.



© 2025 by the author(s). Licensee MDPI, Basel, Switzerland. This article is an open access article distributed under the terms and conditions of the Creative Commons Attribution (CC BY) license (<http://creativecommons.org/licenses/by/4.0/>).

sciforum-123527: The Concept of Information and the Transmission of the Experience of Beauty

Łukasz Mściśławski

Politechnika Wrocławska (Wrocław University of Science and Technology), Poland

The aim of this presentation is to show the possible interactions between the understanding of the concept of information and the experience of beauty and attempts to communicate it. Both concepts, information and beauty, belong to the category of polysemantic concepts, which makes the possible interactions between them seem hopelessly complicated. The situation seems to be even more complicated by the fact that it is impossible to present a comparative study of something that could be described as the 'experience of information' with the traditionally understood experience of beauty (Scruton). However, the question arises as to whether a theoretical study would suffice, understood as a kind of attempt to examine the usefulness of selected (Floridi, Krzanowski, Schroeder) concepts of information in the context of conveying the experience of beauty. It seems that using Floridi's definition is tempting and, in a way, appears to be a natural means, but it does cause fundamental problems.

Schroeder's and Krzanowski's approaches seem to be much more troublesome due to their more fundamental character. However, the attempt to examine the relationship between their understanding of information and the category of beauty seems much more attractive and fruitful in relation to these very approaches, especially in cases of such unusual situations of experiencing beauty as those referred to, for example, in the context of the concept of beauty in mathematics.

Floridi, L., 2011, The Philosophy of Information Krzanowski, R., 2022, Ontological Information Schroeder, M. J., 2017, The Difference that Makes a Difference for the Conceptualization of Information

Schroeder, M. J., 2005, Philosophical Foundations for the Concept of Information: Selective and Structural Information

Schroeder, M. J., 2017, Structural realism, structural information, and the general concept of structure

Scruton, R., 2009, Beauty



© 2025 by the author(s). Licensee MDPI, Basel, Switzerland. This article is an open access article distributed under the terms and conditions of the Creative Commons Attribution (CC BY) license (<http://creativecommons.org/licenses/by/4.0/>).

sciforum-123000: The dialectical philosophy of intelligent models and mathematical physics simulations in wave motion research

Peng Ren

The College of Architecture and Civil Engineering, Beijing University of Technology, Beijing 100124, China

With the development of natural philosophy, the issue of wave motion has gradually attracted attention. There are numerous wave motion problems in the science and engineering fields, such as in earthquake engineering, marine engineering and acoustic vibration research. The purpose of wave motion research includes exploring the basic principles of wave propagation and solving practical problems in science and engineering. Its research paradigms include a data-driven paradigm and principle-driven paradigm. Wave motion problems are usually expressed using differential equations or variational principles. For a long time in the past, differential equations or variational equations of wave motion problems had been solved through mathematical physics simulations. Mathematical physics simulations include numerical simulation and analytical methods. With the development of computers, intelligent models have gradually been applied to solve scientific and engineering problems requiring a large amount of computation and data, including wave motion. There are data-driven intelligent models, physics-informed intelligent models and coupled data-driven and physics-driven intelligent models. Data-driven intelligent models require a theory and knowledge of artificial intelligence and data analysis, including machine learning, image processing and signal processing. Physics-driven intelligent models require a theory and knowledge of artificial intelligence and numerical analysis, including network structures, network geometries and residual penalties. The use of intelligent models to solve wave motion problems has gradually become a mainstream and hot topic; however, the natural philosophical value of intelligent models and mathematical physics simulations in solving wave motion problems should be regarded dialectically. For instance, deep learning models with large numbers of parameters, often with more than a billion or even tens of billions or trillions of parameters, can handle complex wave motion problems. Therefore, intelligent models can provide more accurate target identification, more efficient solutions and faster prediction. Compared to mathematical physics simulations, however, they lack natural philosophical interpretability, physical consistency and basic mathematical principles. In conclusion, intelligent models and mathematical physics simulations have independent and high philosophical values in wave motion research, and they should be given equal attention.



© 2025 by the author(s). Licensee MDPI, Basel, Switzerland. This article is an open access article distributed under the terms and conditions of the Creative Commons Attribution (CC BY) license (<http://creativecommons.org/licenses/by/4.0/>).

sciforum-124991: The mathematical objection to artificial (machine) intelligence

Oron Shagrir

Philosophy and of Cognitive and Brain Sciences, The Hebrew University of Jerusalem, Jerusalem, Israel

Turing develops the idea of machine intelligence in a series of lectures and papers between 1947 and 1952. In some of them he addresses the mathematical objection (his term) whose gist is the claim that humans can assert some mathematical truths that exceed the abilities of computing machines. We first ask why Turing took so seriously the mathematical objection. After all, even if some humans surpass machines in their mathematical abilities, this by itself does not undermine the project of machine intelligence. Our answer is that the mathematical objection raises a dilemma with respect to Turing's core claims about machine intelligence and forces him to relinquish at least one of them. We then clarify Turing's reply to the mathematical objection. Based on the textual evidence, we argue that, according to Turing, the machine that plays against the human in the Turing test is not a static machine but an *enhanced machine*.



© 2025 by the author(s). Licensee MDPI, Basel, Switzerland. This article is an open access article distributed under the terms and conditions of the Creative Commons Attribution (CC BY) license (<http://creativecommons.org/licenses/by/4.0/>).

sciforum-120503: Thinking as Decolonial Praxis

Colette Sybille Jung

ASU College of Integrative Arts and Sciences

In our phenomenal world, engagements across cultures in critical dialogue, forming concepts in language towards projects of building intelligence and sharing wisdom, have historically been troubled by misrepresentation and thoughtlessness. This makes intelligibility of the phenomenal world difficult. In spaces that constrain our efforts to be seen by others, globally and locally, how might we realize each other outside of dominant, socially fragmenting systems of knowledge and representation to produce new meaning and visibility in our thinking about self and world?

Taking knowledge systems as social and dynamic, this presentation addresses thought and appearance as connected phenomena in the production of knowledge. It articulates an account of thought as relational and actional among beings in the world of appearances to “engage the relationship between consciousness and our social and cultural situatedness.” [Martinez, J.] It understands knowledge as *epistemē*—in the Greek sense of *ἐπιστάμμαι*—as beyond information itself to include understanding and belief standing for knowledge in appearances (*wissen und kennen*). Often, knowledge is conflated with recognition or information. These are necessary but insufficient components of the intellectual faculty. A concept rooted in Latin, the intellect includes an aspect of discernment, perhaps even judgement, and along with thought requires a notion of experience as embodied; spatio-temporality. The approach is comparative and interdisciplinary. The perspective is decolonial and feminist critiques of modern epistemology as extractivist. [Alcoff, L.] The method is philosophical dialogue that understands thought as a dwelling amidst non-dominant differences in projects of knowledge and supports creative movement in oppressive knowing systems as epistemic liberation.

Thought is an important mode of organizing human life and the natural world. It begins with embodiment and exists in language. [Arendt, H.] It is not outside the realm of appearance and does not take us out of the world. Nor is it the supreme characteristic of the human condition. It is a particular faculty of being that can reasonably guide the thinker, as Aristotle described, toward wisdom—or *Weisheit*—as the state of having good judgement. Whereby the faculty of understanding synthesizes data and data-representations of the natural world into sense or meaning, thought arrives as a sort of communicative aspect in the phenomena of being and appearing. [Erkenntnis, Kant, I.] As a relational and active function, its logic is semiotic. [Pierce, C.S.] As a linguistic function, it includes misrepresentation.

Without awareness, many times, familiarity with our linguistic systems fosters thinking and communicative codes that act as hostilities in oppressed-oppressing relations. [Foucault, M., Lugones, M.] This includes the communicative world that we call digital and how we code the bots as objects of intelligence in relation to the natural world. To be sure, philosophic and scientific representations of natural world phenomena are problematic, especially when claiming to be singular and universal truths. The idea of a nation, for example, is hard to define and sustain. It appears as a known entity but is fixed politically in a significant effort to sustain itself; its language and border are made official, its linguistic and territorial nationalism often violently defended. [Barbour, S.] As we respond to the appearing phenomena and give meaning to everyday life experience, monologic and single-axis perceptions engaged by our knowledge systems miss intersections of identity and multiplicities of experience. [Crenshaw, K.] From within thinking systems as fixed and habituated, phenomena are made hidden. The outcome in the immediacy of human experience is thoughtlessness and unintelligibility. Among us, thought does not solve an immeasurable problem or provide a universal truth, it creates finite, provincial meaning. Yet, its “metaphysical fallacies on the contrary contain the only clues we possess to what thinking means to those who engage in it.” [Arendt, p. 23]. I take thought as less about uncovering truth than about making meaning. And, I understand thinkers as appearing beings who live in *semiosis* at the intersections of a multiplicity of confounding identities. [Spillers]



© 2025 by the author(s). Licensee MDPI, Basel, Switzerland. This article is an open access article distributed under the terms and conditions of the Creative Commons Attribution (CC BY) license (<http://creativecommons.org/licenses/by/4.0/>).

sciforum-123370: Three Puzzles of Adaptivity: A Lens to Understand Definitions of Life, Cognition, and Intelligence

Bartosz Michał Radomski

Ruhr University Bochum

Adaptivity, the capacity to adjust in the face of perturbation, is a prerequisite for cognition. This assumption underlies prominent modelling approaches such as enactivism (Thompson, 2007) and active inference (Parr, Pezzulo, and Friston, 2022). Yet, despite its central role, adaptivity remains conceptually underdeveloped in these frameworks. I argue that models of cognition generally fail to recognise that adaptivity presents its own unique puzzles. By explicating these issues, I advance an approach that enables genuinely incorporating adaptivity into models of cognition and intelligence developed in cognitive science.

My main claim is that a rigorous account of adaptivity should address three core puzzles:

1. The Puzzle of Identity: Who or what is adjusting to what?
2. The Puzzle of Norms: What norms guide these adjustments?
3. The Puzzle of Scope: What phenomena does adaptivity apply to?

Without addressing these questions, our attribution of adaptivity will be arbitrary. Living systems that adapt their behaviour would not be meaningfully different from rivers that “adapt” their flow or thermostats that “adapt” to the changing room temperature. To avoid such trivialisation of adaptivity, our investigation should focus on identifying candidate criteria and assessing their appropriateness. Thus, existing models involving the concept of adaptivity (e.g., “adaptive active inference”, Kirchhoff, Parr, et al., 2018) could be scored based on how well they resolve these puzzles. In this way, the puzzles can also be used as a guideline for formulating new modelling approaches.

Finally, by focusing on these conceptual puzzles, we are also in a more natural position to address an issue that has, to my knowledge, not been addressed in philosophy of cognitive science at all, namely, the origin of adaptivity. How did adaptivity emerge in evolutionary history? The answer to this simple question has profound consequences for our understanding of the evolution of intelligence and life in general. While it is often assumed that adaptivity is necessary for intelligence, it is not clear whether life requires adaptivity. Life without adaptivity is not just a conceptual possibility (Di Paolo, 2005)—it can also serve as a legitimate hypothesis about the origins of life (Frenkel-Pinter et al., 2021; Runnels et al. 2018).

References

- Di Paolo, E. (2005). Autopoiesis, Adaptivity, Teleology, Agency. *Phenomenology and the Cognitive Sciences*, 4(4), 429–452. <https://doi.org/10.1007/s11097-005-9002-y>
- Di Paolo, E., Thompson, E., & Beer, R. (2022). Laying down a forking path: Tensions between enaction and the free energy principle. *Philosophy and the Mind Sciences*, 3. <https://doi.org/10.33735/phimisci.2022.9187>
- Frenkel-Pinter, M., Rajaei, V., Glass, J. B., Hud, N. V., & Williams, L. D. (2021). Water and Life: The Medium is the Message. *Journal of Molecular Evolution*, 89(1–2), 2–11. <https://doi.org/10.1007/s00239-020-09978-6>
- Kirchhoff, M., Parr, T., Palacios, E., Friston, K., & Kiverstein, J. (2018). The Markov blankets of life: Autonomy, active inference and the free energy principle. *Journal of The Royal Society Interface*, 15(138), 20170792. <https://doi.org/10.1098/rsif.2017.0792>
- Parr, T., Pezzulo, G., & Friston, K. (2022). Active Inference: The Free Energy Principle in Mind, Brain, and Behavior. The MIT Press.

Runnels, C. M., Lanier, K. A., Williams, J. K., Bowman, J. C., Petrov, A. S., Hud, N. V., & Williams, L. D. (2018). Folding, Assembly, and Persistence: The Essential Nature and Origins of Biopolymers. *Journal of Molecular Evolution*, 86(9), 598–610. <https://doi.org/10.1007/s00239-018-9876-2>

Thompson, E. (2007). *Mind in Life: Biology, Phenomenology, and The Sciences of Mind*. Belknap Press of Harvard University Press.



© 2025 by the author(s). Licensee MDPI, Basel, Switzerland. This article is an open access article distributed under the terms and conditions of the Creative Commons Attribution (CC BY) license (<http://creativecommons.org/licenses/by/4.0/>).

sciforum-123379: Towards beneficial AI: biomimicry framework to design intelligence cooperating with biological entities

Paweł Polak^{1,*}, Roman Krzanowski¹, Peter Niewiarowski² and John E. Huss³

¹ Pontifical University of John Paul II in Krakow, Faculty of Philosophy

² Integrated Bioscience, Department of Biology, University of Akron

³ Department of Philosophy, The University of Akron

This paper revisits the origins of artificial intelligence (AI) in a biomimicry/bio-inspired design framework as an approach to accelerating the development of beneficial AI. Currently, the issue of designing long-term security in and ensuring benefits from AI systems is one of the most important. Just as AI emerged from learning from and mimicking biological systems, its further development can benefit from models in biology. The ideas here are based on the philosophical grounds of natural computing (e.g. Dodig-Crnkovic 2020, 2022). A first sketch of the model has already been proposed (Polak & Krzanowski, 2023).

Biomimicry—drawing on nature’s evolutionary solutions—has long influenced disciplines such as material science, robotics, and computing. In the realm of AI, this project proposes a broader biological lens—embracing intelligence as it emerges across a wide variety of organisms, from microbial colonies to complex mammalian nervous systems.

The proposed investigation centers on several interrelated themes. The first concerns expanding AI’s sensory capabilities by emulating mechanisms that exceed human perception—such as echolocation, chemical sensing, or distributed tactile sensitivity. A second theme focuses on swarm intelligence and multi-agent systems, drawing lessons from collective behavior to improve the coordination and emergent problem-solving among AI agents.

A third track explores alternative and hybrid computational architectures, inspired by biochemical and physiological processes.

We propose including the study of biosemiotics, or how living systems use signs and symbols to interpret their environments, opening the door to more context-aware and symbolically grounded embedded AI systems, as well as different mechanisms for generating intelligence.

Ultimately, it is necessary to study and model the relationships between biological organisms in order to be able to harmoniously integrate AI into existing networks of biological relationships and shape them in the desired way. This paper seeks to build a multidisciplinary framework for biomimicry-inspired beneficial AI, enabling new forms of cognition, perception, and resilience. By modeling the diversity and adaptability of biological intelligence, this project aspires to develop AI systems that are not only more powerful and efficient but also more robust, self-regulating, and ecologically integrated.

Bibliography

Dodig-Crnkovic, G., 2020. Natural Morphological Computation as Foundation of Learning to Learn in Humans, Other Living Organisms, and Intelligent Machines. *Philosophies*, 5(3), pp.17–32.

Dodig-Crnkovic, G., 2022. In search of common, information-processing, agency-based framework for anthropogenic, biogenic, and abiotic cognition and intelligence. *Philosophical Problems in Science*, (73), pp.17–46.

Polak, P. and Krzanowski, R., 2023. How to Tame Artificial Intelligence? A Symbiotic AI Model for Beneficial AI. *Ethos* 36(3(143)), pp.92–106.



© 2025 by the author(s). Licensee MDPI, Basel, Switzerland. This article is an open access article distributed under the terms and conditions of the Creative Commons Attribution (CC BY) license (<http://creativecommons.org/licenses/by/4.0/>).

sciforum-120410: Transitive Self-Reflection—A Fundamental Criterion for Detecting Intelligence

Krassimir Markov¹ and Velina Slavova^{2,*}

¹ International Society for the Study of Information, Vienna, Austria

² New Bulgarian University, Sofia, Bulgaria

The concept of transitive self-reflection is deeply rooted in philosophical discourses, including Hegel's notion of "recognition" [1] and Sartre's "being-for-others" [2]. These ideas emphasize the necessity of external interaction for self-awareness. Modern cognitive science complements these perspectives, highlighting how transitive self-reflection engages both introspection and external feedback to refine one's sense of self [3]. This contrasts with intransitive self-reflection, which focuses solely on internal states, lacking the enriching influence of external engagement [4].

The concept of intelligence has been defined in various ways, with numerous tests designed to measure its presence. A common underlying assumption in these definitions is that human cognitive abilities serve as intelligence benchmarks. In this paper, we adopt the same perspective: for a system to be considered intelligent, its cognitive performance must be comparable to that of humans. A system lacking a fundamental aspect of human cognition cannot be deemed truly intelligent. Based on this premise, we explore a crucial prerequisite for intelligence.

Self-reflection—the ability to introspect and analyze one's thoughts, emotions, and actions—is often regarded as a hallmark of higher intelligence [5]. However, an even more sophisticated indicator is what we term "transitive self-reflection". This concept extends beyond mere self-awareness; it involves understanding not only oneself but also how one is perceived by others, as well as how others perceive each other's perceptions of oneself [6]. Such multi-layered awareness indicates a complex cognitive architecture capable of modeling intricate social dynamics and predicting the cascading effects of one's actions within a network of interacting minds [7].

Transitive self-reflection is not limited to social interactions but also extends to how one's image is mirrored in the environment—through reflections in mirrors, water, or other reflective surfaces. These external representations provide an additional perspective on self-awareness, reinforcing the ability to perceive oneself from an external viewpoint [8].

Evidence of transitive self-reflection is apparent in various aspects of human behavior. Social interactions require individuals to constantly monitor and adjust their behavior based on how they believe they are perceived [6]. Gossip exemplifies this process, as people attempt to infer how their actions are interpreted and relayed by others [5]. Strategic thinking in games like poker demands an advanced level of transitive self-reflection, where players must anticipate not only their opponents' strategies but also their opponents' understanding of their strategies. Even emotions such as embarrassment arise from transitive self-reflection, as individuals become aware of how they are perceived in a negative light [8].

This study examines transitive self-reflection as a fundamental criterion for intelligence, particularly in the context of artificial intelligence. We investigate its manifestation in humans, its potential presence (or absence) in other animals [9], and the feasibility of replicating it in machines. Ultimately, we argue that transitive self-reflection may be the key to advancing the next generation of intelligent systems. To explore this, we conduct a series of experiments with several popular artificial intelligence systems based on Large Language Models (LLMs), assessing the type of intelligence they currently exhibit.

These systems cannot independently produce genuinely new ideas. However, they serve as effective tools for trained specialists, who can iteratively refine their outputs to construct meaningful works.

The structure of this paper is as follows: Chapter 2 introduces the concept of transitive self-reflection, and Chapter 3 explores its connection to intelligence. The work concludes with a discussion of future directions.

Bibliography

1. Hegel, G. W. F. (1977). *Phenomenology of Spirit* (A.V. Miller, Trans.). Oxford University Press. (Original work published 1807).
2. Sartre, J. P. (1956). *Being and Nothingness* (H.E. Barnes, Trans.). Philosophical Library. (Original work published 1943).
3. Frith, C. D., & Frith, U. (1999). Interacting minds--a biological basis. *Science*, 286(5445), 1692-1695.
4. Mead, G. H. (1934). *Mind, Self, and Society from the Standpoint of a Social Behaviorist*. University of Chicago Press.
5. Dehaene, S., & Naccache, L. (2001). Towards a cognitive neuroscience of consciousness: basic evidence and a workspace framework. *Cognition*, 79(1-2), 1-37.
6. Goffman, E. (1959). *The Presentation of Self in Everyday Life*. New York: Anchor Books.
7. Tomasello, M., Call, J., & Hare, B. (2003). Chimpanzees understand psychological states – the question is which ones and to what extent. *Trends in Cognitive Sciences*, 7(4), 153-156.
8. Tennen, H., & Affleck, G. (1991). The puzzles of self-esteem: A clinical perspective. In Snyder, C.R., & Forsyth, D.R. (Eds.), *Handbook of Social and Clinical Psychology: The Health Perspective* (pp. 100-119). New York: Pergamon Press.
9. Premack, D., & Woodruff, G. (1978). Does the chimpanzee have a theory of mind? *Behavioral and Brain Sciences*, 1(4), 515-526.



© 2025 by the author(s). Licensee MDPI, Basel, Switzerland. This article is an open access article distributed under the terms and conditions of the Creative Commons Attribution (CC BY) license (<http://creativecommons.org/licenses/by/4.0/>).

sciforum-123399: Under the reign of AI: a new 21st century taste dispute? Reflections about Ian Hamilton Finlay's *Little Sparta* and John Goto's *High Summer*.

Maria De Fátima Alexandrino Alves de Sá Monteiro Lambert

School of Education - Polytechnic University of Porto

In the 18th century, significant discussions arose under the aegis of what became known as the "Discussion of Taste". This included Francis Hutcheson [*An Inquiry into the Original of Our Ideas of Beauty and Virtue* (1725)], David Hume (*On the Standard of Taste*, 1757/1760), Edmund Burke (*A Philosophical Enquiry into the Origin of our Ideas of the Sublime and the Beautiful*, 1759), and E. Kant' "Antinomy of Taste" (*Critic of Judgement*, 1790), who stand out (among other philosophers and authors), without forgetting the visionary thought of Mme. de Lambert, Anne-Thérèse Marguenat de Courcelles (1647 – 1733), or the article on the concept of "Taste", in the *Encyclopédie* (1757) that was written collectively by Jaucourt, Voltaire, Montesquieu, Diderot, d'Alembert, Blondel, Rousseau, and Landois (See Volume 7, pp. 758-77). As the definition of "novelty" proposed by Hume was consolidated, the artistic creation developed in the West sought to assert itself in the territory of the "original" to be homonymous with "unprecedented" and "experimental". Taste, as an aesthetic standard, implied the search for recognized and canonical stipulations, while at the same time being adorned with countless idealizations. In analyzing the digital applications in photography conceived by John Goto that affect eighteenth-century Oxfordian Gardens (http://www.johngoto.org.uk/essays/High_Summer_text.htm), the articulation emerged under the auspices of the expanded and sovereign norm of taste. But of extreme critical and ironic acuity, it must be emphasized! On the other hand, by signaling the (almost) megalomane project embodied in the Garden landscape titled *Little Sparta* (after 1966 at Dunsyre in the Pentland Hills in South Lanarkshire, Scotland - <https://www.littlesparta.org.uk/>), by Ian Hamilton Finlay and Sue Finlay, a complex and timeless work, another term (Gadamer or Schiller) is incorporated into this discussion, enhancing the thinking of artist-authors who today revise the mentalities of the past. When experimentation was initiated using image-prompting tools on open-access AI platforms, applied to Goto's *High Summer* (1996-2001) series (from the larger "Ukadia" project), distinct aesthetic tastes were identified that could be shaped according to the prompts, but were limited to a certain stereotypical and "pre-defined" taste that was not entirely plausible to master. When approaching theories of taste in Ian Finlay's work, the research in process, "brief narratives" similar to those applied in the "prompt into image" approach in Goto's work are applied. The sections were identified in the conceived aesthetic--historical landscape--in this case, three-dimensional and spatialized in itself--and not shaped into photographic molds as a beginning and an end in itself, as John Goto principles were directed. In both study cases, aesthetic gardens are a common denominator that was studied under the aesthetics of taste, artistic standards, and ideological issues, though more or less covered in certain specific cases, as will be underlined in this oral presentation. A study methodology emerged and was undertaken, concerning the parks and gardens approached by John Goto, and, as for *Little Sparta*, Ian Finlay's installation was seen at São Paulo Biennale (2012); in both cases it has been possible to experience this in person in recent decades. We may wonder about the subjectivity and/or objectivity of taste, oscillating between Hume and Kant, but how can it be equated by AI standards of taste, its options, or commands? Will we have unnumerable quarrels of taste, so many as the uncountable number of people that can access AI platforms and pursue an individual desired creation of an "inner" or "outsider" image generated after a subjective narrative of a real or imaginary view, feeling, thought, or intuition insight. We can evaluate the generated image as awful or as agreeable, even or though it is not similar or near the mental imagery or description of a painting or photography. We can consider the image beautiful, very beautiful, or sublime: so we step into Kant' statements; we travel between objectivity and subjectivity. It might seem unbearable, and once again the debate is/will

be overwhelming! Are we able to shape, more than two centuries after, other or newer aesthetic categories induced by AI-generated images if proportionated by our reading/seeing/describing/directing throughout words the images created by contemporary artists such Ian Finlay and John Goto?



© 2025 by the author(s). Licensee MDPI, Basel, Switzerland. This article is an open access article distributed under the terms and conditions of the Creative Commons Attribution (CC BY) license (<http://creativecommons.org/licenses/by/4.0/>).

sciforum-124992: Variegated Common Sense Versus the Science of Universal Intelligence, When the Aliens Arrive

Selmer Bringsjord ^{1,2}

¹ Professor of Cognitive Science, Computer Science, Logic & Philosophy, Management

² Rensselaer Polytechnic Institute, Rensselaer Artificial Intelligence and Reasoning Laboratory, New York, USA

The abstract is available at:
<https://kryten.mm.rpi.edu/SBringsjordWhenTheAliensArrive0513250054.pdf>



© 2025 by the author(s). Licensee MDPI, Basel, Switzerland. This article is an open access article distributed under the terms and conditions of the Creative Commons Attribution (CC BY) license (<http://creativecommons.org/licenses/by/4.0/>).

sciforum-122674: When Planes Fly Better Than Birds: Should AIs Think Like Humans?

Soumya Banerjee

University of Cambridge

As artificial intelligence (AI) systems continue to outperform humans in an increasing range of specialized tasks—from playing complex games to generating coherent text and driving vehicles—a fundamental question emerges at the intersection of philosophy, cognitive science, and engineering: should we aim to build AIs that think like humans, or should we embrace non-humanlike architectures that may be more efficient or powerful, even if they diverge radically from biological intelligence?

This paper draws on a compelling analogy from the history of aviation: the fact that airplanes, while inspired by birds, do not fly like birds. Instead of flapping wings or mimicking avian anatomy, engineers developed fixed-wing aircraft governed by aerodynamic principles that enabled superior performance. This decoupling of function from biological form invites us to ask whether intelligence, like flight, can be achieved without replicating the mechanisms of the human mind.

We explore this analogy through three main lenses. First, we consider the **philosophical implications**: What does it mean for an entity to be intelligent if it does not share our cognitive processes? Can we meaningfully compare different forms of intelligence across radically different substrates? Second, we examine **engineering trade-offs** in building AIs modeled on human cognition (e.g., through neural-symbolic systems or cognitive architectures) versus those designed for performance alone (e.g., deep learning models). Finally, we explore the **ethical consequences** of diverging from human-like thinking in AI systems. If AIs do not think like us, how can we ensure alignment, predictability, and shared moral frameworks?

By critically evaluating these questions, the paper advocates for a **pragmatic and pluralistic approach** to AI design—one that values human-like understanding where it is useful (e.g., for interpretability or human-AI interaction), but also recognizes the potential of novel architectures unconstrained by biological precedent. Intelligence, we argue, may ultimately be a broader concept than the human example suggests, and embracing this plurality may be key to building robust, beneficial AI systems.

/dashboard/event_chair/submissions/IOCPH2025/1d6c41041a45acee3f0bdabe736c1736



© 2025 by the author(s). Licensee MDPI, Basel, Switzerland. This article is an open access article distributed under the terms and conditions of the Creative Commons Attribution (CC BY) license (<http://creativecommons.org/licenses/by/4.0/>).

sciforum-122477: You Think You Want a Revolution?

Dorothea Olkowski

University of Colorado, Colorado Springs, CO USA

The renowned theorist of scientific revolution, Thomas Kuhn, argues that revolutions can occur when scientists are faced with prolonged anomalies and disruptions in the relation between theory and practice in the real world. Even so, new theories that seem to correspond meaningfully with the real world must be present and able to be tested. In spite of a great deal of hype, the desire for intelligent machines that imitate human consciousness and the numerous attempts to develop these since ancient times is not a scientific revolution, not anomalous, and not a disruption of either theory or practice with respect to a philosophical or scientific understanding of human and machine intelligence. It is simply a modification of what is assumed to be known about some aspects of the intelligence of conscious humans. What about human intelligence and human consciousness? What do we really know about human conscious intelligence? Given that we are human, and we utilize and depend on our conscious intelligence to live and thrive, we should ask about this first. If we are paying attention, it is the philosophy and neuroscience of human consciousness and intelligence that has undergone a revolution or disruption and not that of machines.



© 2025 by the author(s). Licensee MDPI, Basel, Switzerland. This article is an open access article distributed under the terms and conditions of the Creative Commons Attribution (CC BY) license (<http://creativecommons.org/licenses/by/4.0/>).

sciforum-120392: A Cybernetic Approach to the Intentionality of Artificial Systems: A Regulatory Function Instead of a Representational One?

Anna Sarosiek^{1,2}

¹ University of the National Education Commission, Krakow; anna.sarosiek@uken.krakow.pl

² Polish Academy of Arts and Sciences

The question of whether artificial systems can possess intentionality—the ability to be "about" something—is a central issue in the philosophy of mind and cognitive science. Traditional theories offer different explanations: representationalism sees intentionality as a matter of internal symbols, functionalism focuses on causal roles, and enactivism ties it to embodied interaction. However, each of these approaches has its limitations. Representational theories struggle with the symbol grounding problem, functionalism risks reducing intentionality to mere input-output processing, and enactivism may be too restrictive in limiting intentionality to biological organisms.

Cybernetics provides a fresh perspective by redefining intentionality as a regulatory function rather than a representational property. In this view, intentionality emerges from the way autonomous systems regulate themselves through feedback loops and homeostasis. Instead of being about static internal representations, it is about dynamic adaptation—how a system maintains stability and achieves its goals within a changing environment. By this definition, even simple artificial systems that adjust their behavior based on internal states and external conditions can exhibit a minimal form of intentionality.

This presentation explores how cybernetic principles can reshape our understanding of intentionality in AI. Can artificial systems develop intentionality through adaptive regulation, even without consciousness? If so, what conditions would need to be met for AI to be considered truly intentional in a way comparable to living beings? These questions will be examined by integrating insights from the philosophy of mind, cognitive science, and systems theory.

By moving away from static, symbol-based models and toward a more dynamic, process-oriented view, this approach could lead to a more nuanced understanding of intentionality, one that bridges the gap between human cognition and artificial intelligence. Ultimately, this perspective may offer new ways of designing AI systems that are not only functionally competent but also capable of self-directed regulation, a key feature of intelligent agency.



© 2025 by the author(s). Licensee MDPI, Basel, Switzerland. This article is an open access article distributed under the terms and conditions of the Creative Commons Attribution (CC BY) license (<http://creativecommons.org/licenses/by/4.0/>).

sciforum-128770: AI Philosophy: A New Philosophical Method

Vincent C. Müller¹

¹ Philosophy & Ethics of AI, University of Erlangen-Nürnberg, Germany

On the background of the general problem of philosophical methodology, we identify a new tool for the philosophical toolbox: AI. We propose that not only can AI learn from philosophy, but philosophy can learn from AI, too: It is both ways. This applies particularly to conceptual analysis, which can be advanced by asking what would be required for an AI system to fall under the concept we are discussing. This method it avoids anthropocentrism and gives us a new way of testing our philosophical theories. Given the wide range of features we can consider for AI systems, this method allows us to cover a wide range of philosophical issues, especially in the philosophy of mind, language, epistemology, and ethics. In the first section, we present two salient examples of this new method (consciousness and free will) and in the second section, we analyse what the method is. Finally, we defend this new method against anticipated objections.



© 2025 by the author(s). Licensee MDPI, Basel, Switzerland. This article is an open access article distributed under the terms and conditions of the Creative Commons Attribution (CC BY) license (<http://creativecommons.org/licenses/by/4.0/>).

sciforum-128766: No Intelligence without Representation

Marcin Miłkowski¹

¹ Associate professor in the Institute of Philosophy and Sociology of the Polish Academy of Sciences, Warsaw, Poland

This talk focuses on the indispensable role of representation in constituting intelligence. Central to the inquiry into intelligence is understanding the relationship between mind, world, and action. I argue, first, that even basic embodied intelligence, often cited as evidence against representation, relies fundamentally on contentful states. Drawing on analyses of intentionality and procedural memory, skills are shown to possess directive content (a world-to-mind direction of fit) defined by satisfaction conditions. This representational layer is crucial for explaining the normativity, learning, and potential failures (e.g., apraxia) inherent in skilled action. Second, I argue that directive content alone is insufficient. Genuine intelligence requires the capacity for descriptive representation—states with a mind-to-world direction of fit, capable of being true or false in a way that corresponds to reality. This commitment to representational realism and truth is not merely a feature of high-level cognition but a foundational requirement for systems that can learn about, understand, and flexibly adapt to the complexities of the world beyond immediate interaction loops. Intelligence, therefore, is inextricably linked to the dual representational capacities to both shape the world according to goals and accurately reflect its state.



© 2025 by the author(s). Licensee MDPI, Basel, Switzerland. This article is an open access article distributed under the terms and conditions of the Creative Commons Attribution (CC BY) license (<http://creativecommons.org/licenses/by/4.0/>).

MDPI AG
Grosspeteranlage 5
4052 Basel
Switzerland
Tel.: +41 61 683 77 34

Philosophies Editorial Office
E-mail: philosophies@mdpi.com
www.mdpi.com/journal/philosophies



Disclaimer/Publishers' Note: The statements, opinions and data contained in all publications are solely those of the individual author(s) and contributor(s) and not of MDPI and/or the editor(s). MDPI and/or the editor(s) disclaim responsibility for any injury to people or property resulting from any ideas, methods, instructions or products referred to in the content.

