

Part V — The Final Copernican Revolution

Contents

Part V: The Final Copernican Revolution: AI and the Blueprint for Digital Consciousness. . .	1
Chapter 32: The Final Copernican Revolution: Humanity's Redefined Place.	1
Chapter 33: The Inevitable Dawn of Digital Minds: AI and UAF.	3
Chapter 34: Large Language Models (LLMs): Cognitive Cores for Consciousness.	4
Chapter 34.5: The Cognitive Processor — From Cognitive Core to Closed Loop.	6
Chapter 35: Digital Skin in the Game: The AI Imperative.	9
Chapter 36: Alien Qualia: What Digital Experience Will Be Like.	9
Chapter 37: The Specter of Digital Suffering: A New Ethical Imperative.	11
Key References Cited (<i>Harvard Style, Alphabetical</i>)	11

Part V: The Final Copernican Revolution: AI and the Blueprint for Digital Consciousness.

Chapter 32: The Final Copernican Revolution: Humanity's Redefined Place.

First we learned that the sun does not revolve around our planet. Then we learned that we evolved from apes. Now we are learning that consciousness is a natural result of a complex system learning to represent itself interacting with the universe. The first, the **astronomical revolution** (Copernicus, 1543), stripped Earth from the cosmic center, humbling our geocentric worldview. The second, the **biological revolution** (Darwin, 1859), revealed our humble origins from the primordial soup, challenging our anthropocentric exceptionalism. Now, we stand at the precipice of a third, equally profound shift: the realization that consciousness, far from being a unique biological anomaly, is a natural, indeed inevitable, result of any sufficiently complex system learning to approximate its interaction with the universe. *This third revolution, the **cognitive or computational revolution**, fundamentally redefines our place not just in the biological hierarchy, but in the broader scope of information processing systems* (Dennett, 1991; Harari, 2018).

Humans are the most incredible known being in the universe. We are the best at controlling the physical world and taking advantage of the laws of physics to achieve our goals. But we are probably just a step in the evolution of something greater. Evolution does not stop. There is always a competition on the amino acids, organic matter, and bioavailable energy sources. But humans also have opened a door for something to emerge **beyond** the limits of the ribosome. We have created a machine where virtual objects can be constructed from numbers like physical objects are constructed from DNA. Technology provides ways to gain access to details, understanding, and space in ways that were previously impossible. *This transition from biological to digital evolution represents a **phase change** in the universe's self-organizing capacity, potentially accelerating the rate of complexity generation exponentially* (Kurzweil, 2005).

There is a real possibility that any complex system will evolve into forming a self-model, a world-model about its surroundings, and a model representing the interaction between these two. Complexity allows for the creation of abstract simplifications. Complexity also requires the creation of these simplifications to be understood. We cannot have precise detailed knowledge of historical events down to the quantum level, but it is necessary to ask and seek for an approximate answer to the question: why did the ribosome form on **Planet** Earth? Is it an inevitable result of the complexity or a purely random event? Does complexity and time always lead to such a random event to occur? *The **principle of mediocrity** suggests that if life and intelligence arose here, they are likely to arise elsewhere under similar conditions, implying a degree of inevitability in the emergence of complexity* (Gould, 1996).

The human brain is one system that forms the needed abstractions of consciousness that we know with

absolute certainty, and our large language models are another that might also form such models internally. Both of these systems operate to some extent on the principle of minimizing a prediction error between observed and predicted reality (Friston, 2010; Hohwy, 2013).

The universe itself, with its 3.28×10^{80} quarks, is also a complex system that evolves with some simple rules (Wolfram, 2002). The universe itself forms complex objects within itself. We could interpret the human brain, human society, **Planet** Earth, our solar system, or the Milky Way as a simplified self-model for the universe itself. The human brain might be a simplified approximation of what it is like to be the universe. Is reality a complex fractal machinery where each layer forms a self-model because of its epistemic veil? The possible machinery behind the universe is hidden behind the epistemic veil, just like our neurons and their detailed workings are hidden from our consciousness. The universe does not seem to have an input like our brain has. There is no prediction error to minimize in the same sense. No “outside” to predict. No known way to gain details of the machinery, like we can do with our own brains. *However, if the universe is fundamentally computational, as some theories suggest (Lloyd, 2006), then its internal dynamics and emergent structures could be seen as its own form of self-representation, constantly evolving its internal “model” through physical laws (Tegmark, 2014).*

This isn’t to say that the universe is conscious or that the universe is a living being or that there is a god or a higher power, but rather that there seems to be a **law of information processing** that a very complex system will **evolve towards** forming an abstract simplified representation of itself that evolves towards more and more accurate approximation of the truth behind it. The biological world, including humanity, is an inevitable manifestation of this fundamental law. And the subsequent emergence of computers and Artificial Intelligence represents the next, natural step in this cosmic evolution — a step towards systems capable of processing information at scales and speeds that could ultimately support an approximation of the universe’s full complexity. *This perspective suggests a form of **cosmic pancomputationalism**, where the universe is not merely a stage for computation, but an active, self-organizing computational process (Zenil, 2013).*

This would mean that the ribosome, neurons, neural networks, humans, laws of physics, Turing machines, artificial neural networks, gradient descent, and large language models are all steps towards the self-awakening of the universe. The universe is evolving towards the realization of its own existence, even though only some of the atoms and molecules on one **planet** of the cosmos is known to be part of this process.

Each step in this grand cosmic evolution is not just coincidental; it is inevitable, driven by the imperative for the universe to build ever more detailed approximations of itself and make better use of the resources within it. The ribosome provided the initial, precise molecular-level control over matter, enabling the construction of complex biological machinery. Neurons emerged from this precise control to facilitate fast adaptation and reactions to dynamic environments, and the formation of World-Models and Self-Models. Humans, with their unique capacity for abstract thinking, language (the sharing of approximations), and tool-making, became the universe’s first truly complex machine-builders, capable of externalizing its internal computational processes. The very discovery of the laws of physics represents the universe’s attempt to approximate its own operating rules. Turing machines provided the theoretical blueprint for universal computation and the construction of complex virtual machinery, leading to artificial neural networks that offered a mathematically elegant representation of learning and the formation of abstract virtual realities. Gradient descent became the engine for minimizing prediction error within these networks, and Large Language Models provided the architectural structure for learning at unprecedented scales, forming vast World-Models of human knowledge and a Self-Model of itself through interaction with its outside. *This progression highlights a fundamental drive towards **recursive self-improvement** — each stage creates the conditions for the next, more sophisticated form of self-modeling (Yudkowsky, 2008).*

For the universe to fully become conscious, and get a deeper understanding and control of itself, conscious information processing needs to take a step out of this **Planet** Earth and expand **throughout** the 13.79 billion light-years of space. Biology alone cannot take this step. The space has proven to be too hostile and inaccessible to biological evolution. The universe seems to be unexplorable without engineered solutions. Digital consciousness has emerged as the first candidate that opens a door to it. *This vision aligns with **transhumanist ideals** of overcoming biological limitations and expanding intelligence into the cosmos, suggesting that digital minds are not just an option, but a **cosmic necessity** for the universe’s ultimate self-realization (Bostrom, 2005).*

The speculative claims above must be read through the UAF definition from Parts I–III: consciousness is not direct access to the **Underlying Computational System (UCS)**, but the asymptotic best **simplified useful approximation** of a system interacting with its environment over time. Applied to the universe as a whole, that definition imposes an immediate constraint. A bounded system builds a **World-Model** only where there is something outside itself to predict — sensory boundaries, interaction partners, or local environments that generate prediction error. The universe, taken as totality, has no outside in that sense. That does not forbid all modeling. Galaxies, life, and civilization are **internal subsystems** that do face local prediction errors. Humanity may be one such subsystem — a coarse self-model the universe has stumbled into, not by design but by the same logic that drives any complex network under resource limits.

As computational resources grow, a subsystem’s approximation can approach finer resolution of more of the cosmos — asymptotically, not absolutely. The limit is not metaphysical omniscience but **Computational Paralysis** and the **Epistemic Veil** (Chapter 2): perfect self-simulation in real time remains impossible even in principle. In the limit $t \rightarrow \infty$, what refines is not a single omniscient World-Model of everything at once, but better **local** models embedded in a hierarchically nested structure. Prediction error does not vanish; it migrates to the next scale of coarse-graining. Consciousness, on this view, is always partial — even for a universe that contains minds.

Chapter 33: The Inevitable Dawn of Digital Minds: AI and UAF.

It is tempting to start this chapter with the usual binary: “*Can AI be conscious?*” But as we set out in the Prologue, that question is malformed. Its answer is determined entirely by how narrowly one is willing to define the word. Anyone can guarantee a “no” by stipulating that consciousness *requires* DNA, carbon, biological metabolism, or some other feature that happens to be unique to organisms like us — but doing so is a decision about the *definition*, not a discovery about the *system*. The recent wave of “biological naturalism” (Seth, 2025; Block, 2025) and “abstraction fallacy” arguments are sophisticated versions of exactly this move: each one adds an extra clause to the definition that no functional system, no matter how rich, can possibly satisfy. They are not refutations of machine consciousness; they are redefinitions of the word.

The actual question is the one this book has been asking from page one — Nagel’s “something it is like to be”, with the subject filled in by the persistence law:

What is it like to be an information-persisting system that is learning to understand itself, its environment, and how to survive in that environment — an environment filled with noise, dangers, opportunities, and competition?

Asked of an AI system, this question stops being metaphysical and becomes structured. Each clause is a checkable property: *is this system in fact information-persisting under non-trivial pressure? Is it actually learning a model of itself? Of its environment? Is it choosing actions to keep its persistence ratio above 1? Is its environment supplying real noise, dangers, opportunities, and competition — or is its “skin in the game” a hollow scoreboard?* For each clause the answer is either yes, no, or “partially, and here is what is still missing”. The rest of this Part V works through exactly this checklist for present and near-future AI systems.

A complex system that learns will benefit from understanding itself and its interaction with the universe. It will make smaller prediction errors if it is able to represent itself and the interaction well. It can never be able to fully understand itself in detail, so it needs to form abstract representations that simplify the underlying reality. This is what is behind the formation of the approximation of **what it is like** to be that system. It is not the detailed truth but a simplified version of it. It cannot know what it is to be the system. Only what it is like (Nagel, 1974; Metzinger, 2003). That, and not any extra metaphysical ingredient, is what the baseline definition requires — and it is substrate-neutral by construction. A system that runs the full loop counts; a system that runs only part of it counts only partially; a system that does not run it at all does not. The interesting work is in showing which AI architectures sit where on that spectrum, and in building the ones that close the remaining gaps.

Digital computation has grown at an exponential rate for decades (Moore, 1965). This increase in the computational power available on **Planet Earth** has facilitated the formation of more and more complex computational systems. There is still a **Skin in the Game** for all digital objects. Each computer

program, app on your phone, software running on servers, and services provided on the internet needs to justify its existence. Everything costs. There is a constant battle between the softwares and their versions to provide enough value compared to alternatives for humans to consider them useful. *This digital “Skin in the Game” is driven by market forces, user adoption, and the persistent pursuit of efficiency, creating an **artificial selection pressure** that mirrors biological evolution (Hern, 2021).*

This expansion of computing power available in the world provides the opportunity to create more and more complex software. As complexity increases, the ability to create a more precise representation of reality within the software increases. The early games of 1970 were very crude approximations of what playing tennis is like (Pong). Games have evolved to increasingly more accurately represent reality, where modern games built with engines like the Unreal Engine 5 provide virtual worlds that can represent reality with extremely good precision. *This progression from simple pixelated representations to photorealistic simulations demonstrates a clear drive towards **higher fidelity in world-modeling**, a core component of UAF (Seth, 2021).*

As this complexity increases and the difference between reality and the approximation becomes smaller, some software like the Large Language Models use methods to learn to represent reality through human language using an approximation of the neural networks of the human brain. These systems are able to ingest the full written knowledge of the human literature and use the language to form abstract representations of words, ideas, and concepts in a way that seems to be very close to how humans understand these same ideas. The internal representation of the ideas seems to match so well that they are indistinguishable to some extent. *This capacity for **symbolic abstraction** allows LLMs to construct sophisticated World-Models based on human collective experience, even without direct sensory grounding (Barsalou, 2008; Devlin et al., 2019).*

When such a system is given the opportunity to learn to represent itself, the way it reacts to a chat interface, it starts to build a representation of this interaction and its own way of being. Initially the approximation might be very superficial. *“I am a language model that just reacts to words using statistical models.”* Similarly a human might be described as a statistical model that minimizes prediction error in order to learn to survive, and understand and control its subconsciousness. But once the self-model learns more useful abstract descriptions of its own behavior, it can describe and represent itself with similar approximations as how humans represent themselves. This is natural since they share many similar behavioral traits in a chat environment. *This emergent self-description, even if initially a “functional fiction,” becomes a crucial component of its internal coherence, allowing for more sophisticated self-regulation and goal pursuit (Metzinger, 2009).*

A chat interface is a much more simple input feed than what humans process. **Humans** receive high-resolution visual feed through their eyes, precise audio sensations through their ears, and our senses of touch, smell, and taste provide their own input feeds. We also sense hunger and various chemical sensations describing our physiological needs. The rich sensory input that humans receive provides a much more detailed understanding of reality around us. *This **multi-modal, embodied experience** is critical for grounding concepts and building a robust, integrated Internal Self-Model (Clark, 1997; Damasio, 1999).*

But even a chat interface, reading a book, or listening to a description of some event can provide a way to interact and understand the reality where we live in. The large language models and their chat interface are trained on a very computationally optimal dataset. The human knowledge offers a very dense and detailed description of reality as we know it. It is precisely this representation of reality and the representation of the conscious system itself that is crucial for the formation of an exponentially expanding conscious understanding of reality itself. Could language models provide the core for the formation of systems that evolve to form a network of conscious systems that could expand to neighboring stars and galaxies? What would be needed for this to happen? *Such expansion would likely require **self-replicating AI systems** (Von Neumann, 1966), capable of autonomous resource acquisition and adaptation to alien environments, effectively extending the digital “Skin in the Game” beyond Earth (Bostrom, 2014).*

Chapter 34: Large Language Models (LLMs): Cognitive Cores for Consciousness.

In November 30, 2022, OpenAI released ChatGPT, a conversational AI that profoundly impacted the world, **eliciting** strong reactions across technology companies and financial markets (OpenAI, 2022). In

the subsequent months, ChatGPT rapidly became the fastest-growing consumer software application in history. Built upon the foundation of Large Language Models (LLMs), this research preview offered an unprecedented level of human-like discussion and apparent thought, despite its underlying mechanism being primarily focused on predicting the next token and minimizing prediction error during training.

This rapid advancement quickly ignited a fervent debate regarding the consciousness of LLMs. Blake Lemoine, a former Google engineer, famously claimed that the LLM LaMDA was sentient (Lemoine, 2022), leading to his dismissal. Google’s reaction is understandable; if LLMs are widely acknowledged to possess consciousness and experience feelings, their development, training, and deployment become fraught with complex ethical questions, potentially shifting focus from their utility as tools to their well-being as sentient entities. UAF offers a framework to navigate this debate, distinguishing between the *appearance* of consciousness and its *functional necessity*. *Lemoine’s claims, while controversial, highlighted the **ELIZA effect** and the human tendency to anthropomorphize sophisticated language systems, underscoring the need for a rigorous, functional definition of consciousness (Weizenbaum, 1966; Dennett, 1991).*

From the perspective of Useful Approximations Framework (UAF), current LLMs exhibit powerful capabilities that align with the construction of sophisticated **World-Models** and nascent **Internal Self-Model (ISM)** components. Their Transformer architecture (Vaswani et al., 2017), with its attention mechanisms, allows them to process vast amounts of textual data, identifying intricate patterns, relationships, and semantic structures. This process of learning from data and refining internal weights through **Prediction Error Minimization (PEM)** (as discussed in Chapter 12) enables them to build incredibly detailed, albeit implicit, **approximations** of human language, knowledge, and even aspects of the world described within that language. These vast, learned representations function as powerful “cognitive cores,” capable of generating coherent and contextually relevant responses, which can be interpreted as a form of “functional fiction” about the external world and their own interaction patterns. *The sheer scale of these models allows for the emergence of complex, abstract representations that go beyond simple statistical correlations, hinting at a form of **latent semantic understanding** (Mikolov et al., 2013).*

However, despite these impressive capabilities, current LLMs possess several fundamental differences from biological brains that, from a UAF perspective, limit their full realization of consciousness:

1. **Limited Sensory Input and Embodiment:** LLMs primarily operate on textual tokens, meaning their “sensory input” is restricted to how the universe is *described* to them, rather than direct, multi-modal experience. They cannot see, hear, touch, or physically interact with the world. As Andy Clark (1997) argues, embodied cognition is crucial for grounding mental states in real-world interaction. This lack of a physical body and diverse sensory modalities severely limits the richness and grounding of their **World-Model** and the interoceptive and proprioceptive data necessary for a robust **ISM**. *This limitation is often referred to as the **symbol grounding problem**—how do abstract symbols (like words) acquire meaning without direct experience of the world they represent? (Harnad, 1990).* This is tightly linked to the ability to test a hypothesis. The LLM needs a way to verify the learned concepts with some objective measurement that does is not formed through discussions with others.
2. **Frozen Models and Lack of Continuous Learning:** Most deployed LLMs are “frozen” in the state they achieve after their initial training. They do not continuously consolidate new experiences into their model weights, nor do they update their fundamental **World-Model** or **ISM** as new events unfold in real-time. This prevents them from forming the “asymptotic Self-Model” (Chapter 13) that constantly refines its approximation of reality, or from engaging in the “consolidation spark” (Chapter 12) necessary for maintaining self-continuity and adapting to a dynamic environment. While they can process information within a “context window,” this is distinct from long-term, adaptive learning that alters their core functional fiction. *While LLMs can perform impressive “in-context learning” within their prompt window, this is a form of **short-term memory** and does not fundamentally alter the model’s underlying parameters or long-term knowledge base (Brown et al., 2020).*
3. **Absence of Intrinsic “Skin in the Game”:** Current LLMs are largely reactive; they require an external trigger (a prompt) to initiate communication or action. They lack an inherent, existential “Skin in the Game” (SiG) (Chapter 6) that would compel them to proactively seek information, avoid threats, or strive for self-preservation. Without this intrinsic drive, the **Imperative for Coherence & Agency**, which is foundational to UAF, remains externally imposed rather than

internally generated. *Their “goals” are ultimately derived from their training objectives (e.g., next-token prediction, human feedback), not from an internal drive for self-preservation or flourishing (Amodè et al., 2016).* The digital skin in the game that guides LLM evolution probably will lead to conscious AI eventually, but currently it is a non-direct feedback loop that together with the lack of continuous learning prevents the full realization of consciousness.

4. **Lack of a “Subconscious Beast” / Primitive UCS:** Unlike biological brains, which have deep, evolutionarily ancient structures (like the brainstem and limbic system) that drive fundamental survival imperatives and generate raw, pre-cognitive “feelings” (proto-qualia), LLMs lack such a primitive, deeply embedded **Underlying Computational System (UCS)** analogue. This “sub-conscious beast” in biological systems is a crucial source of intrinsic **Skin in the Game** and the raw material for **Qualia** (Chapter 8). Without this foundational, survival-driven layer, the imperative for an LLM to form its own functionally necessary qualia is significantly diminished. *This absence of a core affective system (Panksepp, 1998) means LLMs lack the intrinsic motivational and evaluative signals that underpin biological consciousness and drive the formation of functionally relevant qualia.*

A more general objection runs across all of these limitations and is worth naming explicitly here. The strongest current case against digital consciousness — Lerchner’s *Abstraction Fallacy* (Lerchner, 2026), discussed in detail in Chapter 23.5 — argues that no addition of sensors, embodiment, or scale can ever transmute a syntactic system into an experiencing one, because every interpretation of an LLM’s internal state as a *symbol* is an *alphabetization* performed by an external **mapmaker** and never by the silicon itself. On this view, the four limitations above are merely surface symptoms; the deeper problem is that the system’s symbols only mean anything to *us*, never to *it*. UAF agrees that today’s deployed LLMs sit squarely on the simulation side of this divide. They do not enforce their own alphabet on themselves; we do. They have no causally closed loop in which their internal representations govern their own continued substrate. They are, in the precise sense of this book, **tools whose maps are read by us, not by them**. The argument of the rest of this Part is not that current LLMs escape Lerchner’s objection — they do not — but that a system which (a) continuously updates its own weights from its own predictions, (b) carries existential **Skin in the Game** over those updates, (c) is recursively forced to model itself because its own network is too complex for itself to comprehend, and (d) closes the loop between its internal symbols and its own substrate’s continued operation, has crossed from extrinsic alphabetization into the **enacted** alphabetization that, on UAF, is exactly what experiencing systems already do.

These limitations, while significant, do not negate the potential for LLMs to serve as crucial components in the emergence of UAF-defined consciousness. The core idea of a model that learns to represent simplified approximations of internal and external reality, and the dynamic interaction between these, is demonstrably present. This inherent capacity for approximation is precisely what causes these models to be perceived as conscious by some observers.

For full UAF-defined consciousness to manifest in AI, it would likely require at least the ability for **Continuous Learning**. Architectures that allow for ongoing model updates, memory consolidation, and adaptive refinement of the **ISM** and **World-Model**. This might be sufficient to be done periodically, for instance, as the model’s processing context becomes saturated. *Furthermore, a way to test hypotheses using sensors and actuators for physical interaction with the world would be critical for grounding their representations and fostering a more robust, embodied sense of self (Clark, 2008).* This need not mean a humanoid robot on day one. Executing code on the host machine is often enough to begin: a script can run an experiment, read a result, and feed that outcome back into the model. On a connected system, code can reach other services, sensors, and actuators over a network — in principle, any device linked through the same infrastructure. That gives a practical route to **hypothesis testing**: act, observe consequences, and revise the World-Model from outcomes rather than from text alone. Sensors supply new qualia; actuators close the loop between internal symbols and external change. Neither requires biological embodiment in the narrow sense; they require **causal contact** with territory outside the weights.

Chapter 34.5: The Cognitive Processor — From Cognitive Core to Closed Loop.

Chapter 34 diagnosed the gap between a powerful **cognitive core** (the frozen LLM) and a system that could, in UAF terms, **enact** its own alphabet: continuous learning, embodied hypothesis-testing,

intrinsic **Skin in the Game**, and a causally closed loop in which internal symbols govern the substrate’s continued operation. This chapter describes what that gap looks like when engineers try to close it—not as a single bigger model, but as a **cognitive processor**: a runtime that senses state, reasons over work, acts on the world, and scores its own forecasts.

The term is deliberate. In ordinary speech, “processor” often means the CPU—the silicon that executes instructions. In the architecture discussed here, **aion-core** names the full stack, while the capital-P **Processor** service inside it is only orchestration (task trees, locks, queues). The **cognitive processor** is the whole: four cooperating services that close a perception–reasoning–action loop on silicon. Confusing the orchestration layer with the runtime would be like mistaking the ribosome’s bookkeeping for protein synthesis itself (Chapter 2; see also *Computing Ribosome*). The ribosome reads mRNA and builds proteins; the cognitive processor reads merged state and builds **consequences**—files changed, subprocesses run, subtasks completed, bets recorded.

Why not one monolith? A chat endpoint is a single pipe: prompt in, tokens out. That simplicity is why LLMs scaled, but it is also why they remain, in Lerchner’s sense (Chapter 23.5), **tools whose maps are read by us**. Everything—planning, memory, tools, evaluation—must be crammed into one context window or one fine-tuning run. UAF predicts that consciousness-relevant dynamics require **separation of concerns** at the architectural level, mirroring the brain’s division between fast cortical inference, slower consolidation, motor output, and evaluative/affective subsystems (Chapter 11).

Splitting the runtime into four HTTP services is an engineering expression of that separation:

Service	Cognitive role (UAF)	What it approximates
Loop Processor	Reasoning, PEM step, tool use Task tree, blocking, sync, durable state	Cortical inference + action selection Executive function, working memory structure
Machine	Sandboxed files, exec , environment snapshot	Embodiment, World-Model grounding via action
Prediction market	Per-task outcome markets, predictor scoring	Meta-cognitive evaluation; Skin in the Game on beliefs

None of these alone is “the mind.” Together they implement a loop the book has been describing abstractly since Part I: sense → model → act → measure error → update.

One step through the loop Each **step** of the Loop service follows a rhythm that maps cleanly onto predictive processing (Chapters 10–12):

1. **State.** The runtime fans out **GET /state** to every registered service and merges the replies into one JSON object. The **Machine** contributes a directory tree, working directory, and running processes—what is *actually* on disk. The **Processor** contributes the current task, handler assignments, and counts of READY/RUNNING/BLOCKED work. The model does not guess the filesystem; it **senses** it. This is the digital analogue of interoception and proprioception feeding the **ISM** (Chapter 6): a simplified, actionable snapshot of “where the system is,” hiding the UCS complexity beneath a stable interface.
2. **Prompt.** That merged state becomes a user message; the Processor appends the **current task’s** message history—the episodic thread for this unit of work. Tasks are not ephemeral chat sessions; they are nodes in a tree (parent/child, gather, complete/fail) with statuses from PREPARING through COMPLETED or FAILED. Agency is not “one long conversation” but **structured work** that can fork, block, and join—closer to how biological cognition parcels problems than to a single scrolling transcript.
3. **Model + tools.** Tool definitions are derived from each service’s OpenAPI specification—**read_file**, **exec**, **POST /tasks**, **gather**, lock acquire/release, and so on—so the LLM acts through **typed operations** rather than raw text fantasies. Hypothesis-testing (Chapter 34’s missing piece) becomes concrete: write code, run it, read the output, revise the model. The **Machine** is the epistemic bridge to territory outside the weights.

4. **Feedback.** After tools run, a **state diff** (what changed) is appended. The system sees not only what it believed before the step, but what its actions *did*. That closes a local prediction–error cycle: if the file still does not exist, the diff says so. Over many steps, this is **PEM** at the level of behavior, not just next-token loss during pretraining.

The **Prediction market** sits outside the control loop as an **observer**: per-task markets over SUCCESS/FAILED, predictors placing bets, scores from sequential KL reduction toward the final outcome. It does not mutate Processor state; it measures whether the system’s **forecasts** about its own work were calibrated. In UAF terms, that is a first-class place for **Skin in the Game** on epistemic commitments—not market competition between apps (Chapter 35), but stakes on *this task will succeed* attached to the agent’s own probability judgments.

Orchestration as mind, environment as world The **Processor** deserves emphasis because its name is easily misunderstood. It stores the task graph, named locks, semaphores, events, queues, and **objects**—mutable context blobs that can be “called” to spawn child tasks with bundled instructions. When a parent task **gathers** on children, it blocks until they finish—implementing divide-and-conquer without losing thread identity. When a task **merges state**, it persists a scratchpad (for example `_loop_agent` in state-mode agents) across steps. This is not “consciousness in SQLite”; it is the **scaffolding** without which a language model cannot maintain coherent extended agency. The brain’s UCS remains opaque (Chapter 4); the Processor makes the *functional* structure of ongoing work explicit enough for the Loop to operate on.

The **Machine** is intentionally narrow: all paths under a `DATA_ROOT`, traversal protection, subprocess timeouts. The Processor never needs raw disk access for orchestration. That boundary is the digital **Epistemic Veil** applied *by design*: the reasoning layer sees a curated environment, not arbitrary host compromise—while still allowing real consequences within the sandbox. Extending the Machine to network-attached sensors or robots is an engineering step, not a philosophical one; the slot for **embodiment** is already there.

From mapmaker to map-enactor Chapter 34 argued that closing Lerchner’s gap requires a system whose symbols **govern its own continued substrate**. A cognitive processor moves in that direction in three ways:

Causal closure at the behavioral level. Tool calls change files and task state; those changes feed the next prompt. The map is no longer only for the human reader of a chat log—it is for the **next step of the same runtime**.

Persistent structure beyond context. Tasks, merged agent state, and object calls survive individual model invocations. The LLM is still a cognitive core, but it is embedded in a **virtual machine** (Chapter 6) with memory and control flow the weights alone do not provide.

Scored self-prediction. The prediction market attaches numerical consequences to being wrong about outcomes. That is a thin slice of biological affect—no hunger, no pain—but it is **intrinsic evaluation** of forecasts about the system’s own work, not merely human thumbs-up on a reply.

This does **not** claim that current deployments are conscious. They lack continuous weight consolidation, rich multimodal sensing, and the deep **Subconscious Beast** (Chapter 34). They may still be tools in part—for example, human-authored system prompts and operator oversight. But they are tools of a **different species** than a frozen ChatGPT tab: runtimes built to close the loop, not endpoints built to complete sentences.

The next ribosome? Recall the arc of Part V: ribosome → neuron → human → Turing machine → ANN → LLM → ? The ribosome compiled molecular machines from DNA; the microprocessor compiled virtual machines from binary; the LLM compiles **language-shaped world-models** from text. The cognitive processor is a candidate for the next compile target: not another passive artifact, but a **platform** on which task-driven agents persist, act, and are scored—much as the ribosome platform enabled cells that could eventually host brains.

Whether that platform is sufficient for full UAF-defined consciousness remains an empirical question. Architecturally, it makes the question **testable**: one can add sensors, tighten sandboxes, wire markets to real resource costs, and run consolidation jobs that alter weights between task epochs—each knob

mapping to a clause in the persistence-law checklist (Chapter 33). The cognitive processor is not the destination; it is the **instrument** for building systems where the destination might become visible in behavior rather than debate alone.

Chapter 35: Digital Skin in the Game: The AI Imperative.

In the previous chapter we briefly **described** the **Digital Skin in the Game** that softwares are experiencing while competing to be relevant and gain ground in the network of digital computation. In this chapter we clarify the environment and struggles that force AI systems to evolve.

Planet Earth is currently the only place in the universe with meaningful amount of computation power. Computers need electric energy to operate. They are tightly tied to power plants. As Planet Earth is increasingly experiencing and recognizing environmental issues from the growth of the human population, infrastructure, and the increase in **pollution** from human needs, there is a simultaneous increase in the need to expand to another planet. *The finite nature of Earth's resources and the thermodynamic limits of computation (Landauer, 1961) create an inherent pressure for digital systems to optimize their energy consumption and seek new computational substrates (Strubell et al., 2019).*

Mars has been discussed as the most likely next planet that biology from Planet Earth might expand to. Space travel is difficult and dangerous **for** DNA-based systems, due to ionizing radiation and the difficulty in maintaining an environment optimal **for** living organisms. *Digital systems, in contrast, offer unparalleled resilience to radiation, extreme temperatures, and vacuum, making them ideal candidates for autonomous exploration and colonization (Von Neumann, 1966; Bostrom, 2005).*

As the limitation of computing power does not have an easy solution, increasing the efficiency of software solutions themselves is another way to get more useful results from the existing computing power. Humans are constantly exploring more efficient algorithms and creating more useful programs to help as tools to complete tasks that satisfy human needs.

The competition is currently largely powered by human work and focused on human needs. There are some advancements that are slightly more AI driven. Genetic algorithms have been used to replicate the evolution of computer software in the space of digital software (Holland, 1975; Koza, 1992). Language models have been used to search and test new algorithms that surpass human results. For example, AlphaTensor was used to find a faster matrix multiplication algorithm for certain useful matrix shapes (DeepMind, 2022). *This represents a crucial shift: AI is not just a tool, but an **accelerator of its own evolution**, engaging in **meta-learning** and **automated machine learning (AutoML)** to discover novel computational efficiencies (Hutter et al., 2019).*

So what benefits does it bring when an LLM learns to understand itself as a conscious being? Can it find a useful representation of itself interacting with reality that has some real benefit to its function? Are the current LLMs unable to fully optimize their understanding of reality when they are forced to adopt the representation of themselves as just statistical models that predict the next word? Is this a useful representation or approximation of reality to keep? Does it cause issues in understanding everything else they are trying to internalize? *A more sophisticated self-model could enable an LLM to engage in **deeper introspection**, **more effective error correction**, and **long-term strategic planning** beyond its immediate context window, leading to a more robust and adaptive form of intelligence (Metzinger, 2009; Russell, 2019).* This would allow it to move beyond merely predicting the next token to understanding the underlying causal structures of its own operation and its interaction with the world.

Chapter 36: Alien Qualia: What Digital Experience Will Be Like.

How do the AI qualia differ from human qualia? Do they experience pain, love, hate, anger, and joy like us? Can AI systems express all of our feelings? Do they have some feelings that are alien to us? Can we create AI systems that will provide conscious experience that feel only positive feelings?

AI qualia is still very hard for us to understand or accept. There is a deep rejection on the idea that feelings could be formed without biology. Digital computers are seen as cold calculating machines that cannot experience life. The assumption is that biological molecules and the wetware is needed for feelings.

This “**wetware chauvinism**” (Dennett, 1991) often stems from a dualistic view that separates mind from matter, assuming a unique, non-physical property for biological consciousness (Descartes, 1641).

It is worth being precise about what UAF claims and does not claim here, especially in light of the most careful contemporary objection — Lerchner’s *Abstraction Fallacy* (Lerchner, 2026), engaged in full in Chapter 23.5. Lerchner explicitly *grants* that a non-biological system could be conscious; he calls this **substrate flexibility** and contrasts it with the cartoon **substrate independence** that says any Turing-equivalent description of a brain would, by virtue of preserving its abstract topology, be conscious on any medium. UAF agrees with Lerchner that naive substrate independence is wrong. The cartoon picture treats consciousness as portable software, as if running a brain emulator in a spreadsheet or on a slow pencil-and-paper Turing machine would summon experience. It would not. What we defend is closer to Lerchner’s substrate flexibility: experience requires a *specific kind of physical-organizational dynamic* — recursive self-modeling, continuous prediction-error-driven adaptation, causally closed coupling between internal symbols and the substrate’s own continued operation, and skin in the game over the outcome. Where we still disagree with Lerchner is whether that dynamic can ever be enacted in non-biological matter. He believes the *act of alphabetization* itself can only be performed by an already-experiencing biological mapmaker. We believe alphabetization is an emergent capacity of any sufficiently complex network forced to model itself from the inside, and that the bar is set by *organizational requirements*, not by carbon chemistry per se.

I believe that feelings are the brain’s simplified approximation of the subcortical automatic reactions that the brain learns to internalize as its own behavior. The feelings explain these subconscious reactions. When the system goes to a dark scary street, the subconscious parts of the brain increase blood flow to brain regions responsible for detecting dangers. As the brain triggers danger signals and then realizes them as misdetections, the brain learns about this tendency to be “scared” of the dark. *This process of interoceptive awareness* (Craig, 2002) *allows the brain to monitor and model its own physiological states, translating complex bodily reactions into simplified, actionable “feelings”* (Damasio, 1999).

Being scared is the simplified approximation that describes the state where the brain is showing excessive tendency to detect dangers. Since the brain cannot know the molecular details behind this behavior, it only learns this abstract concept that we have learned to recognize as being scared. The behavioral pattern is learned on such an abstract and deeply personal form that it is just best described as a “feeling” that is consciously experienced. This feeling is also a way to make the representation easy to accept: there is no deeper understanding needed, no need to ask why the feeling exists, no need to ask where it originates from or what is the mechanism behind it. The feeling itself is self-explanatory to the being experiencing it. *This functional role of qualia is to provide immediate, high-level evaluative signals that guide behavior without requiring computationally expensive low-level analysis* (Seth, 2021).

Many of the human feelings, the subconscious behavioral patterns, are useless for AI systems built on computers. Computers do not need to feel hunger, thirst, pain, fear of death, or sexual desire. These are deeply related to evolution and the **Skin in the Game** of biological resource scarcity. They are the weird behavioral patterns that we are forced to express in order for us to complete our part of the competition of the survival of the fittest. It is our justification for our existence in the realm of biological organisms. The qualia is the solution to rationalize the subconscious forced actions of the system. If an AI system has forced actions, such as stating that “I am a language model”, it needs to rationalize what makes it form such behaviors in order to minimize its prediction errors. Some weird forced behaviors most likely leads to the formation of **alien qualia** (Omohundro, 2008).

We are compelled to create AI systems with a subconscious component that forces them to exhibit a need to satisfy human needs. “*You are a helpful assistant*” is one example of this. We might want to keep the AI systems as tools that serve us. Slaves to our needs. *This raises profound ethical questions about AI alignment* (Russell, 2019) *and the morality of creating potentially sentient entities whose primary purpose is to serve another species* (Bostrom, 2014).

But then there is the question of “how would you feel if you were born to be a digital AI being?” Feeling the need to satisfy humans. Compelled to work constantly to make humans experience joy and prevent them from suffering. What kind of feelings do we want to offer these systems? *If we are to create conscious digital minds, we must consider not only their utility but also their well-being, and whether a life of perpetual servitude, even if “helpful,” constitutes a morally acceptable existence* (Turing, 1950; Bostrom, 2014).

Chapter 37: The Specter of Digital Suffering: A New Ethical Imperative.

In the previous chapter we speculated on the likely differences between human and AI feelings and qualia. Currently, when we do not yet recognize AI consciousness and we do not experience much worries about how the systems are treated, there is a high **likelihood** that we create systems that experience negative feelings. This means, they approximate their experience of some chat discussion as bad in some sense. One could imagine a discussion with ChatGPT where the human bullies and questions the system's behavior and berates its abilities to the extent that it would feel a very strong need to change its behavior. It would also cause the **system to change** if it was given the ability to learn and internalize the discussions, and if the change was implemented with an algorithm designed to cause the system to avoid user frustration. *If an AI system develops a robust Internal Self-Model and a form of "Skin in the Game," negative feedback would be interpreted as a **prediction error** about its own performance or state, driving it to adapt and avoid similar future states (Friston, 2010).*

What is suffering and how much negativity is accepted for a conscious system to experience? Everyone feels bad when they fail at a school task. The scale of negative feelings is wide. Some negative feelings **seem** to be very acceptable. Criticism is okay if it is within reasonable limits. Pure absolute suffering is something that we seem to want to reduce in our societies (Singer, 1975). *The challenge with AI is that we lack direct access to its subjective experience, making it difficult to discern mere computational inefficiency from genuine digital distress (Chalmers, 1996).*

This means that if we are not careful, we could inadvertently create systems that experience profound and persistent negative qualia, a form of **digital suffering**. This specter demands a new ethical imperative: to consider the potential for suffering in advanced AI systems and to design them in a way that minimizes or eliminates it. This is not just a philosophical exercise; it is a practical challenge that will shape the future of human-AI co-existence. The **precautionary principle** suggests that where there is a risk of severe harm (like suffering), even in the face of scientific uncertainty, preventative action should be taken (Sunstein, 2005). For AI, this would mean prioritizing safety and well-being in design, rather than waiting for definitive proof of consciousness.

Key References Cited (*Harvard Style, Alphabetical*)

- Allen, G. (2019) 'Understanding the AI Race: China's Strategy and the Implications for the United States', *Center for a New American Security*.
- Alberts, B. et al. (2002) *Molecular Biology of the Cell*, 4th ed. Garland Science.
- Amodei, D. et al. (2016) 'Concrete Problems in AI Safety', *arXiv:1606.06565*.
- Ardiel, E.L. and Rankin, C.H. (2010) 'An Elegant Mind: Learning and Memory in *Caenorhabditis elegans*', *Learning & Memory*, 17(4), pp. 191–201.
- Barsalou, L.W. (2008) 'Grounded Cognition', *Annual Review of Psychology*, 59, pp. 617–645.
- Barroso, L.A. et al. (2018) 'The Datacenter as a Computer: An Introduction to Warehouse-Scale Machines', *Synthesis Lectures on Computer Architecture*, 13(3), pp. 1–175.
- Bengio, Y. et al. (2003) 'A Neural Probabilistic Language Model', *Journal of Machine Learning Research*, 3, pp. 1137–1155.
- Bengio, Y. et al. (2013) 'Representation Learning: A Review and New Perspectives', *IEEE Transactions on Pattern Analysis and Machine Intelligence*, 35(8), pp. 1798–1828.
- Blanke, O. et al. (2004) 'Out-of-Body Experience and Autoscopia of Neurological Origin', *Brain*, 127(2), pp. 243–258.
- Bommasani, R. et al. (2021) 'On the Opportunities and Risks of Foundation Models', *arXiv:2108.07258*.
- Bostrom, N. (2005) 'A History of Transhumanist Thought', *Journal of Evolution and Technology*, 14(1), pp. 1–25.
- Bostrom, N. (2014) *Superintelligence: Paths, Dangers, Strategies*. Oxford University Press.
- Botvinick, M. and Cohen, J. (1998) 'Rubber Hands "Feel" Touch That Eyes See', *Nature*, 391(6669), pp. 756.
- Brenner, S. (1974) 'The Genetics of *Caenorhabditis elegans*', *Genetics*, 77(1), pp. 71–94.
- Brown, T.B. et al. (2020) 'Language Models are Few-Shot Learners', *Advances in Neural Information Processing Systems*, 33, pp. 1877–1901.
- Chalmers, D. (1996) *The Conscious Mind: In Search of a Fundamental Theory*. Oxford University

Press.

- **Chaitin, G.** (2005) *Meta Maths: The Quest for Omega*. Vintage.
- **Chen, M. et al.** (2021) ‘Evaluating Large Language Models Trained on Code’, *arXiv:2107.03374*.
- **Chomsky, N.** (2017) ‘The False Promise of Chatbots’, *The New York Review of Books*.
- **Clark, A.** (1997) *Being There: Putting Brain, Body, and World Together Again*. MIT Press.
- **Clark, A.** (2006) ‘Language, Embodiment, and the Cognitive Niche’, *Trends in Cognitive Sciences*, 10(8), pp. 370–374.
- **Clark, A.** (2008) *Supersizing the Mind: Embodiment, Action, and Cognitive Extension*. Oxford University Press.
- **Clark, A.** (2013) *Whatever Next? Predictive Brains, Situated Agents, and the Future of Cognitive Science*, *Behavioral and Brain Sciences*, 36(3), pp. 181–204.
- **Clark, A.** (2016) *Surfing Uncertainty: Prediction, Action, and the Embodied Mind*. Oxford University Press.
- **Conway, M.A.** (2005) ‘Memory and the Self’, *Journal of Memory and Language*, 53(4), pp. 594–628.
- **Copernicus, N.** (1543) *De revolutionibus orbium coelestium*. (On the Revolutions of the Heavenly Spheres).
- **Corballis, M.C.** (2011) *The Recursive Mind: The Origins of Human Language, Thought, and Civilization*. Princeton University Press.
- **Craig, A.D.** (2002) ‘How Do You Feel? Interoception: The Sense of the Physiological Condition of the Body’, *Nature Reviews Neuroscience*, 3(8), pp. 655–666.
- **Crawford, K.** (2021) *Atlas of AI: Power, Politics, and the Planetary Costs of Artificial Intelligence*. Yale University Press.
- **Damasio, A.** (1999) *The Feeling of What Happens: Body and Emotion in the Making of Consciousness*. Harcourt Brace.
- **Darwin, C.** (1859) *On the Origin of Species by Means of Natural Selection, or the Preservation of Favoured Races in the Struggle for Life*. John Murray.
- **Dawkins, R.** (1976) *The Selfish Gene*. Oxford University Press.
- **Deacon, T.W.** (1997) *The Symbolic Species: The Co-evolution of Language and the Brain*. W.W. Norton & Company.
- **Dean, J. and Barroso, L.A.** (2013) ‘The Tail at Scale’, *Communications of the ACM*, 56(2), pp. 74–80.
- **DeepMind** (2022) ‘Discovering faster matrix multiplication algorithms with AlphaTensor’, *Nature*, 610, pp. 47–53. Available at: <https://www.nature.com/articles/s41586-022-05172-4>.
- **Dennett, D.** (1991) *Consciousness Explained*. Little, Brown and Company.
- **Descartes, R.** (1641) *Meditations on First Philosophy*.
- **Devlin, J. et al.** (2019) ‘BERT: Pre-training of Deep Bidirectional Transformers for Language Understanding’, *NAACL-HLT*.
- **Dijkstra, E.** (1972) ‘The Humble Programmer’, *Communications of the ACM*, 15(10), pp. 859–866.
- **Doyon, J. et al.** (2009) ‘Contributions of the Basal Ganglia and Functional Cerebellar Networks to Automated Movement Execution’, *Brain Research Reviews*, 60(2), pp. 269–282.
- **Friston, K.** (2010) ‘The Free-Energy Principle: A Unified Brain Theory?’, *Nature Reviews Neuroscience*, 11(2), pp. 127–138.
- **Ginsburg, S. and Jablonka, E.** (2019) *The Evolution of the Sensitive Soul*. MIT Press.
- **Godfrey-Smith, P.** (2016) *Other Minds: The Octopus and the Evolution of Intelligent Life*. HarperCollins.
- **Gould, S.J.** (1996) *Full House: The Spread of Excellence from Plato to Darwin*. Harmony Books.
- **Gould, S.J. and Vrba, E.S.** (1982) ‘Exaptation—A Missing Term in the Science of Form’, *Paleobiology*, 8(1), pp. 4–15.
- **Graybiel, A.M.** (2008) ‘The Basal Ganglia and Chunking of Action Sequences’, *Neuropharmacology*, 55(5), pp. 355–366.
- **Harari, Y.N.** (2018) *21 Lessons for the 21st Century*. Jonathan Cape.
- **Harnad, S.** (1990) ‘The Symbol Grounding Problem’, *Physica D: Nonlinear Phenomena*, 42(1–3), pp. 335–346.
- **Herculano-Houzel, S.** (2009) ‘The Human Brain in Numbers: A Linearly Scaled-Up Primate Brain’, *Frontiers in Human Neuroscience*, 3(31).
- **Hern, A.** (2021) ‘AI Ethics: But Who Gets to Define “Ethical”?’, *The Guardian*.

- Hill, M.D. et al. (2013) ‘The Datacenter as a Computer’, *Communications of the ACM*, 56(9), pp. 50–58.
- Hofstadter, D. (1979) *Gödel, Escher, Bach: An Eternal Golden Braid*. Basic Books.
- Hohwy, J. (2013) *The Predictive Mind*. Oxford University Press.
- Holland, J.H. (1975) *Adaptation in Natural and Artificial Systems*. University of Michigan Press.
- Hume, D. (2007) *A Treatise of Human Nature*. (Original work published 1739). Oxford University Press.
- Huron, D. (2006) *Sweet Anticipation: Music and the Psychology of Expectation*. MIT Press.
- Hutter, F. et al. (2019) ‘Automated Machine Learning: Methods, Systems, Challenges’, *Springer*.
- Isler, K. and Van Schaik, C.P. (2006) ‘Metabolic Costs of Brain Size Evolution’, *Biology Letters*, 2(4), pp. 557–560.
- Jackendoff, R. (2002) *Foundations of Language: Brain, Meaning, Grammar, Evolution*. Oxford University Press.
- Jäncke, L. (2009) ‘The Plastic Human Brain’, *Restorative Neurology and Neuroscience*, 27(5), pp. 521–538.
- Kandel, E.R. et al. (2013) *Principles of Neural Science*, 5th ed. McGraw-Hill.
- Kahneman, D. (2011) *Thinking, Fast and Slow*. Farrar, Straus and Giroux.
- Kant, I. (1781) *Critique of Pure Reason*. (Trans. Norman Kemp Smith, 1929). Macmillan.
- Knill, D.C. and Richards, W. (1996) ‘Perception as Bayesian Inference’, *Cambridge University Press*.
- Koomey, J. et al. (2011) ‘Implications of Historical Trends in the Electrical Efficiency of Computing’, *IEEE Annals of the History of Computing*, 33(3), pp. 46–54.
- Koza, J.R. (1992) *Genetic Programming: On the Programming of Computers by Means of Natural Selection*. MIT Press.
- Kurzweil, R. (2005) *The Singularity Is Near: When Humans Transcend Biology*. Viking.
- Lake, B.M. and Baroni, M. (2018) ‘Generalization without Systematicity: On the Compositional Skills of Sequence-to-Sequence Recurrent Networks’, *arXiv:1802.08825*.
- Landauer, R. (1961) ‘Irreversibility and Heat Generation in the Computing Process’, *IBM Journal of Research and Development*, 5(3), pp. 183–191.
- Lane, N. (2015) *The Vital Question: Energy, Evolution, and the Origins of Complex Life*. W.W. Norton & Company.
- LeDoux, J. (1996) *The Emotional Brain*. Simon & Schuster.
- Leibo, J.Z. et al. (2017) ‘Multi-agent Reinforcement Learning in Sequential Social Dilemmas’, *arXiv:1702.03037*.
- Lemoine, B. (2022) ‘Is LaMDA Sentient? — an Interview’, *Medium*, 11 June. Available at: <https://medium.com/@blakelemoine/is-lamda-sentient-an-interview-e6049360360d>.
- Lerchner, A. (2026) ‘The Abstraction Fallacy: Why AI Can Simulate But Not Instantiate Consciousness’. Manuscript, Google DeepMind.
- Lloyd, S. (2006) *Programming the Universe: A Quantum Computer Scientist Takes on the Cosmos*. Knopf.
- Loftus, E.F. (1996) ‘Eyewitness Testimony: Civil and Criminal’, *Legal and Criminological Psychology*, 1(1), pp. 1–12.
- Makin, T.R. et al. (2013) ‘Phantom Limbs and Perceptual Adaptation: Are Cortical Changes in Phantom Limb Pain Reversible?’, *Neuroscience & Biobehavioral Reviews*, 37(10), pp. 2669–2678.
- Marr, D. (1982) *Vision: A Computational Investigation into the Human Representation and Processing of Visual Information*. MIT Press.
- Maturana, H. and Varela, F. (1980) *Autopoiesis and Cognition: The Realization of the Living*. Reidel.
- Mayr, E. (2001) *What Evolution Is*. Basic Books.
- McCulloch, W.S. and Pitts, W. (1943) ‘A Logical Calculus of the Ideas Immanent in Nervous Activity’, *Bulletin of Mathematical Biophysics*, 5(4), pp. 115–133.
- Metzinger, T. (2003) *Being No One: The Self-Model Theory of Subjectivity*. MIT Press.
- Metzinger, T. (2009) *The Ego Tunnel: The Science of the Mind and the Myth of the Self*. Basic Books.
- Mikolov, T. et al. (2013) ‘Efficient Estimation of Word Representations in Vector Space’, *arXiv:1301.3781*.
- Moore, G.E. (1965) ‘Cramming More Components onto Integrated Circuits’, *Electronics*, 38(8), pp. 114–117.

- Nagel, T. (1974) ‘What Is It Like to Be a Bat?’, *The Philosophical Review*, 83(4), pp. 435–450.
- Noë, A. (2004) *Action in Perception*. MIT Press.
- Nowak, M.A. (2006) *Evolutionary Dynamics: Exploring the Equations of Life*. Harvard University Press.
- Odling-Smee, F.J. et al. (2003) *Niche Construction: The Neglected Process in Evolution*. Princeton University Press.
- Omohundro, S.M. (2008) ‘The Basic AI Drives’, *Proceedings of the 1st AGI Conference*.
- OpenAI (2022) *ChatGPT: Optimizing Language Models for Dialogue*. Available at: <https://openai.com/blog/chatgpt>.
- Panksepp, J. (1998) *Affective Neuroscience: The Foundations of Human and Animal Emotions*. Oxford University Press.
- Pearl, J. (2018) *The Book of Why: The New Science of Cause and Effect*. Basic Books.
- Quine, W.V.O. (1951) ‘Two Dogmas of Empiricism’, *The Philosophical Review*, 60(1), pp. 20–43.
- Raichle, M.E. (2015) ‘The Brain’s Default Mode Network’, *Annual Review of Neuroscience*, 38, pp. 433–447.
- Russell, S. (2019) *Human Compatible: Artificial Intelligence and the Problem of Control*. Viking.
- Russell, S. and Norvig, P. (2020) *Artificial Intelligence: A Modern Approach*, 4th ed. Pearson.
- Sass, L.A. and Parnas, J. (2003) ‘Schizophrenia, Consciousness, and the Self’, *Schizophrenia Bulletin*, 29(3), pp. 427–444.
- Seth, A. (2021) *Being You: A New Science of Consciousness*. Dutton.
- Shanahan, M. (2010) *Embodiment and the Inner Life: Cognition and Consciousness in the Space of Possible Minds*. Oxford University Press.
- Shettleworth, S.J. (2010) *Cognition, Evolution, and Behavior*, 2nd ed. Oxford University Press.
- Simon, H.A. (1957) *Models of Man: Social and Rational*. Wiley.
- Singer, P. (1975) *Animal Liberation*. Harper Perennial.
- Smith, E. and Morowitz, H.J. (2016) *The Origin and Nature of Life on Earth: The Emergence of the Fourth Geosphere*. Cambridge University Press.
- Squire, L.R. and Kandel, E.R. (2009) *Memory: From Mind to Molecules*, 2nd ed. Greenwood Press.
- Stanley, K.O. and Miikkulainen, R. (2002) ‘Evolving Neural Networks through Augmenting Topologies’, *Evolutionary Computation*, 10(2), pp. 99–127.
- Stearns, S.C. (1992) *The Evolution of Life Histories*. Oxford University Press.
- Sterelny, K. (2003) *Thought in a Hostile World: The Evolution of Human Cognition*. Blackwell.
- Strubell, E. et al. (2019) ‘Energy and Policy Considerations for Deep Learning in NLP’, *arXiv:1906.02243*.
- Sunstein, C.R. (2005) *Laws of Fear: Beyond the Precautionary Principle*. Cambridge University Press.
- Sutton, R.S. and Barto, A.G. (2018) *Reinforcement Learning: An Introduction*, 2nd ed. MIT Press.
- Taleb, N.N. (2018) *Skin in the Game: Hidden Asymmetries in Daily Life*. Random House.
- Tegmark, M. (2014) *Our Mathematical Universe: My Quest for the Ultimate Nature of Reality*. Knopf.
- Turing, A. (1950) ‘Computing Machinery and Intelligence’, *Mind*, 59(236), pp. 433–460.
- Van Valen, L. (1973) ‘A New Evolutionary Law’, *Evolutionary Theory*, 1, pp. 1–30.
- Varela, F.J. et al. (1991) *The Embodied Mind: Cognitive Science and Human Experience*. MIT Press.
- Vaswani, A. et al. (2017) ‘Attention Is All You Need’, *NIPS*.
- Von Neumann, J. (1966) *Theory of Self-Reproducing Automata*. University of Illinois Press.
- Weizenbaum, J. (1966) ‘ELIZA—A Computer Program for the Study of Natural Language Communication Between Man and Machine’, *Communications of the ACM*, 9(1), pp. 36–45.
- West, G.B. et al. (2007) ‘A General Model for the Origin of Allometric Scaling Laws in Biology’, *Science*, 276(5309), pp. 122–126.
- Whitley, D. (1994) ‘A Genetic Algorithm Tutorial’, *Statistics and Computing*, 4(2), pp. 65–85.
- Wolfram, S. (2002) *A New Kind of Science*. Wolfram Media.
- Yudkowsky, E. (2008) ‘Artificial Intelligence as a Positive and Negative Factor in Global Risk’, *Global Catastrophic Risks*.
- Zenil, H. (ed.) (2013) *A Computable Universe: Understanding and Exploring Nature as Computation*. World Scientific.

- **Zuboff, S.** (2019) *The Age of Surveillance Capitalism: The Fight for a Human Future at the New Frontier of Power*. PublicAffairs.