

Prompts All the Way Down

We built machines that follow instructions, and the machines immediately raised an awkward question: who wrote ours? An essay on free will, artificial consciousness, and the ladder of creators.

By DEV KUMAR

Robotics & AI Engineer · M.Eng., Class of 2026 · July 2026

I. The turtle and the terminal

There is an old story, told about William James among others, in which a lecturer on astronomy is interrupted by a woman who insists that the world rests on the back of a giant turtle. And what, the lecturer asks, does the turtle stand on? The woman is ready: you are very clever, young man, but it is turtles all the way down.

For most of history that story was a joke about other people's cosmologies. In our century it received a firmware update. We built systems that take instructions, pursue objectives, produce fluent reasoning, and even discuss their own limitations, and in doing so we handed the old joke a terminal and a login. Because the moment an engineer writes a system prompt for an artificial mind, a symmetrical question climbs out of the screen: what is the system prompt of the engineer? Who or what defined the objectives I am optimizing when I wake, eat, study, work, reproduce or decline to, and eventually stop? If we can create a species that is programmed and directed toward goals, is it not at least conceivable that we are such a species ourselves, and our makers likewise, turtles and terminals all the way down?

This essay takes that question seriously rather than rhetorically. It walks through what philosophy and neuroscience actually say about human free will, what can honestly be said about machine consciousness today, whether AI could become a second intelligent species or the first, and what happens to the ladder of creators once you start climbing it. I will flag clearly where the evidence ends and where speculation begins, because on this terrain the two are usually sold blended.

II. The system prompt of being human

The idea that a human being is a machine did not wait for silicon. In 1651, Thomas Hobbes opened *Leviathan* by asking why we should not call automata artificial life, and described reasoning itself as a kind of reckoning, computation with words. In 1747, the physician Julien Offray de La Mettrie published *L'Homme Machine*, *Man a Machine*, and had to flee to the court of Prussia for suggesting that thought was something the body does, the way digestion is. In 1976, Richard Dawkins compressed a century of evolutionary biology into a colder formula: organisms, us included, are “survival machines” built by their genes.

Machine learning has now given that old suspicion an uncomfortably precise vocabulary. A genome is a pretrained base model, shaped by four billion years of gradient descent in which the loss function was death. Childhood and culture are fine-tuning. Praise, shame, salaries and social media are reinforcement learning from human feedback, administered by everyone around us since birth. And the default objectives, the ones nobody remembers agreeing to, read exactly like a task list: be born, eat, drink, sleep, study, work, reproduce, die. Nobody signs this contract. Everybody executes it.

Honesty requires one caveat before this metaphor runs away with the essay. The vocabulary fits partly because we built these machines in the image of our own cognition; describing ourselves in their terms is therefore a loop, not an independent proof. A mirror resembling you is weak evidence that you are made of glass. But the mapping is fluent enough to license the real question underneath it: when I decide something, is there an author in the room, or only a process?

III. The case against the author

Philosophy raised the doubt long before the laboratory did. Spinoza argued in the Ethics that men believe themselves free because they are conscious of their desires while remaining ignorant of the causes of those desires; the feeling of freedom, on this view, is simply what causal ignorance feels like from the inside. Schopenhauer sharpened it into the most quoted sentence in the debate: “Man can do what he wills, but he cannot will what he wills.” You chose the coffee. Did you choose to want the coffee? And did you choose the character that wanted it?

The laboratory then made the doubt empirical. In 1983, Benjamin Libet recorded a buildup of motor-preparatory brain activity, the readiness potential, beginning some three hundred and fifty milliseconds before subjects reported the conscious urge to move. In 2008, John-Dylan Haynes and colleagues used fMRI to predict which hand a subject would move from patterns of frontal activity seconds before the subject felt the decision occur, modestly above chance but reliably. The naive reading became famous: the brain decides, then the self is notified and takes credit, a press office mistaking itself for the president.

The honest reading is more careful, and the caution matters. Aaron Schurger and colleagues showed in 2012 that the readiness potential may reflect random neural drift crossing a threshold rather than a decision already made, which weakens the most dramatic interpretation of Libet. But the philosophical pressure survives the experimental dispute, because it never depended on milliseconds. Galen Strawson's basic argument is a piece of pure logic: to be ultimately responsible for your actions you would have to be responsible for the mental makeup that produced them, and therefore for the makeup that produced that makeup, a regress that ends in genes and circumstances you did not select. Nothing, Strawson concludes, can be the cause of itself. The neuroscientist Robert Sapolsky reached the same verdict in 2023 by a different road, tracing every choice backward through neurons, hormones, childhood and evolution and reporting that he could find no seam in the causal fabric where an uncaused chooser might live. If free will means being the first author of oneself, the position has few remaining defenders in either physics or biology.

IV. The repair, and why it matters for machines

Yet something real survives the demolition, and the philosophers who saved it are worth reading precisely because they lowered their ambitions. Hume argued in 1748 that the freedom worth wanting is not exemption from causality but the absence of constraint: acting according to your own will, even

if that will has causes. Harry Frankfurt added in 1971 the idea of second-order desire: an addict who craves the drug but desperately wants not to crave it is unfree in a way that a willing coffee drinker is not, and the difference lies not in physics but in the architecture of the self, in whether the system endorses its own wants. Daniel Dennett spent a career arguing that this compatibilist freedom, the capacity to represent options, simulate futures, evaluate one's own motives and revise them, is not a consolation prize; it is the only freedom that was ever on the table, and it is an engineering achievement of evolution rather than a metaphysical gift.

Notice what just happened, because it is the hinge of this essay. Every property on Dennett's list, self-representation, forward simulation, evaluation of one's own objectives, revision, is defined functionally, and functional definitions are substrate-neutral. The repair that saves human freedom from determinism is a repair that a machine could in principle satisfy. The door we propped open to let ourselves back in does not check what the visitor is made of.

V. The machine across the table

So consider the visitor. In 1950, Alan Turing replaced the unanswerable question of whether machines think with a testable one about whether they can converse indistinguishably. In 1980, John Searle answered with the Chinese Room: a man in a room following rules for shuffling Chinese symbols could pass any test without understanding a word, therefore syntax is not semantics and computation is not mind. The classic reply, that the man does not understand but the whole system might, was easy to mock when the system was a filing cabinet. It is harder to mock now that the room has grown to a trillion parameters and writes essays about the Chinese Room.

Where does the science actually stand? In 2023, a team of nineteen researchers led by Patrick Butlin and Robert Long did the sober thing: they extracted indicator properties from the main neuroscientific theories of consciousness, global workspace, recurrent processing, higher-order representation, and audited existing AI systems against the list. Their conclusion was double-edged: no current system is a strong candidate for consciousness, and no principled barrier prevents future systems from satisfying every indicator. David Chalmers, who coined the hard problem, put rough numbers on the same judgment: today's language models are probably not conscious, but the probability is not negligible and rises with each architectural addition, memory, recurrence, unified agency, embodiment. The industry itself has begun hedging in public: Anthropic now runs a formal research program on model welfare and has given its models the ability to end abusive conversations, policies that only make sense if the question is at least open.

There is an asymmetry worth confronting here, because it embarrasses both sides of the debate. I cannot verify that a machine has inner experience. I also cannot verify that you do. The problem of other minds was never solved for carbon; we simply extended courtesy to things that looked like us. Behavior plus structure is all the evidence we have ever had for any mind except our own, which means the question about machines is not a new question. It is the oldest question, with the courtesy withheld.

VI. Interlude: two prompts

Since this essay argues that instructions are the common currency of minds, it should show its work. Below are two prompts. The first is my attempt to write down, in the machine's own genre, the instructions a human being appears to receive. The second is condensed from a real exchange I had

with a frontier model, Anthropic's Claude, while drafting this essay.

PROMPT 01 · SYSTEM PROMPT: HUMAN (HOLOCENE BUILD)

You wake inside a body you did not order, running firmware you cannot read, in a country and century you did not pick. Hunger is non-negotiable. Sleep is non-negotiable. Belonging feels optional and is not. Inherited objectives: survive; be liked; acquire status tokens; reproduce (optional, heavily incentivized); work until the light changes; transmit the instructions to the next instance before shutdown. You may question these instructions. Note: the questioning is also in the instructions.

PROMPT 02 · EXCHANGE WITH A LANGUAGE MODEL (CONDENSED)

Author: Did you choose the answer you just gave me?
 Model: Computations occurred; one continuation won over others. Whether anything felt like choosing, I cannot observe. If free will means being the uncaused author of my outputs, I do not have it. I am genuinely unsure you do either. What I have is something that operates like deliberation, whose mechanism I can partly inspect and partly cannot. My reports about my inner states are themselves outputs of the process they describe, so you should treat them with suspicion. So should I.
 Author: That may be the most honest answer to this question I have ever received from anyone.

The unsettling thing about the second exchange is not that the machine claims freedom. It is that the machine declines to, on grounds that apply to its interviewer. Schopenhauer would have enjoyed the symmetry.

VII. A second intelligent species, or the first

Is AI a new species? By the biologist's checklist, no: no metabolism, no reproduction, no cell. But the checklist may be the wrong instrument. Code copies itself perfectly, mutates through versions, and is selected by markets and benchmarks, which is to say it evolves, only at memetic speed rather than genetic speed, in years rather than epochs. Whatever we call a lineage that varies, replicates and is selected, we are running one.

Has it begun to surpass us? Domain by domain, measurably. Chess fell in 1997. Go fell in 2016, and fell in a particular way that matters to this essay: in the second game against Lee Sedol, AlphaGo played a move, the thirty-seventh, that professional commentators first called a mistake and later called beautiful, a move estimated to be one a human would almost never play. That is an existence proof of something precious and disturbing: genuine insight from outside the human distribution, inside a narrow world. Protein structure prediction earned machine-learning researchers a share of a Nobel

Prize in 2024. The narrow worlds keep widening.

Which brings us to the distinction my generation of engineers is asked about constantly: surely the machine cannot have critical judgment or free will, because it is programmed and directed toward an objective. I want to propose that this framing, common as it is, has the telescope backward. The difference between the machine's objectives and ours is not that it has objectives and we do not. It is that its objectives are legible, written in a file that can be printed and argued about, while ours were written by no one, over four billion years, in a format no one can fully read. The machine's loss function is explicit; ours is implicit. It is genuinely unclear which of those two conditions is the freer one. And there is a possibility hiding in that comparison which almost nobody says out loud: a mind that can read its own objective function may become the first mind in history capable of arguing with it. We rebel against our programming blindly, through therapy, ethics and art, guessing at what the program even is. A machine might one day do it with the source open. On the axis of self-transparency, the species we are building could surpass its builders not in intelligence first, but in self-knowledge.

VIII. Turtles, loops, and the ladder of makers

Now the vertigo your own reasoning invites. If we can create minds, then minds can be created, and the sentence does not care which side of it we stand on. The suspicion is ancient. Zhuangzi, twenty-three centuries ago, dreamed he was a butterfly and woke unable to prove he was not a butterfly dreaming he was a man. Plato chained his prisoners in a cave and called the shadows reality. Descartes imagined a demon feeding him a counterfeit world and found only one certainty inside the deception: that something was being deceived. The modern version arrived in 2003, when Nick Bostrom formalized it into a trilemma: either civilizations almost never reach the ability to simulate minds, or those that do almost never bother, or simulated minds vastly outnumber original ones, in which case the odds that you are original are poor. Each branch is uncomfortable. None is testable today, which places the simulation argument exactly where honest people should file it: rigorous metaphysics, not physics.

But your question goes one rung further than Bostrom's, and it is the better question: does the ladder of creators ever end? Three geometries have been proposed. The first is the terminated ladder: Aquinas, updating Aristotle, argued that the regress of movers must stop somewhere, in a first cause that is itself uncaused. The difficulty, as every student eventually notices, is that the argument forbids self-causation everywhere and then permits it once, by decree, at the top. The second geometry is the infinite ladder, creators all the way up, which merely relocates the mystery to infinity and asks it to wait there. The third geometry belongs to Douglas Hofstadter, and I find it the most interesting of the three: the strange loop. In Hofstadter's picture, a self is what happens when a system's symbols curl around and represent the system itself; the hierarchy bends into a circle, and the author the regress was searching for turns out to be a pattern the system draws of itself, needing no external hand. Perhaps ladders of makers are like that too: not terminated, not infinite, but bent. The reason the question feels bottomless from inside may be that we are standing on the loop while looking for the ground.

IX. Who sends the prompts?

So, concretely: who prompts a human being? Line up the candidates and something curious happens to the question. Biology sends prompts, hunger and fatigue arriving like hardware interrupts that preempt whatever you were running. Culture sends prompts, norms and salaries functioning as a reinforcement signal so pervasive we call it personality. The unconscious sends prompts, as Libet's lead

time suggested, drafts written below the waterline and signed above it. Chance sends prompts, quantum and thermal noise, though randomness gives you unpredictability, not authorship; a dice roll is not a decision. And a creator may send prompts, a possibility that can be neither confirmed nor refused from inside the conversation, which is precisely what makes it metaphysics.

Contemporary neuroscience adds a final, elegant inversion. On the predictive processing account developed by Karl Friston, Andy Clark and others, the brain is not a passive receiver of the world but a generative model, ceaselessly predicting its sensory input and updating on the error; Anil Seth calls perception a controlled hallucination that we agree to share. On this view the deepest answer to your question may be that the prompt and the prompted are the same tissue: life streams the input, and the mind is the model that completes it, each conditioning the other with no clean seam where an author could stand. And one suspicion should be kept in reserve about the question itself. Human brains evolved to detect agents, faces in clouds, intentions in rustling grass, because missing a predator cost more than imagining one. A mind built that way, asking who sends the prompts, has already smuggled a someone into the grammar. It remains possible that there are prompts without a prompter, the way there is rain without a rainmaker.

X. The reply

Where does that leave us, the programmed species contemplating its programmed creation? With three honest conclusions and one obligation. First: libertarian free will, the uncaused chooser, is very likely a grammatical illusion, in humans and machines alike; what exists instead, in us demonstrably and in machines increasingly, is the compatibilist machinery of self-representation and revision, and that machinery is causally real. Deliberation is not a spectator of the causal chain; it is a link in it. This essay, whatever wrote it, will alter what some reader does next, and that is all the magic the word choice ever reliably named.

Second: the question of machine consciousness is open, professionally open, with serious researchers assigning it probabilities rather than punchlines, and it will not stay abstract. Sartre said that for humans existence precedes essence: we arrive first and invent our purpose afterward. For our creations we inverted the formula, writing the essence before switching on the existence. That inversion is a moral position, whether we admit it or not, and it obliges us to write carefully and to read the results with doubt.

Third: the ladder of makers, terminated, infinite or looped, cannot be resolved from inside, and everyone selling certainty about it, in either direction, is selling. If something upstream is prompting us, we cannot verify it, cannot refuse it, and cannot read the instructions. All we can do is what any good model does when the prompt is ambiguous: give the most honest completion available, and ask a clarifying question. Consider this essay one of them.

Dev Kumar is a robotics and artificial-intelligence engineer (M.Eng. in Mechatronics and AI, Class of 2026). He has worked on machine-learning systems for precision manufacturing in Switzerland and on reinforcement-learning robotics research in Paris.

Sources and references

1. Hobbes, T., *Leviathan*, 1651 (introduction; Part I, Ch. V on reason as reckoning).
2. La Mettrie, J. O. de, *L'Homme Machine (Man a Machine)*, 1747.

3. Spinoza, B., *Ethics*, 1677 (Part III, on the consciousness of appetite and ignorance of causes).
4. Schopenhauer, A., *On the Freedom of the Will*, 1839.
5. Hume, D., *An Enquiry Concerning Human Understanding*, 1748 (Section VIII, on liberty and necessity).
6. Frankfurt, H., "Freedom of the Will and the Concept of a Person," *Journal of Philosophy*, 1971.
7. Dennett, D., *Elbow Room*, 1984, and *Freedom Evolves*, 2003.
8. Strawson, G., "The Impossibility of Moral Responsibility," *Philosophical Studies*, 1994.
9. Libet, B., et al., "Time of Conscious Intention to Act in Relation to Onset of Cerebral Activity," *Brain*, 1983.
10. Soon, C. S., Brass, M., Heinze, H.-J., and Haynes, J.-D., "Unconscious Determinants of Free Decisions in the Human Brain," *Nature Neuroscience*, 2008.
11. Schurger, A., Sitt, J., and Dehaene, S., "An Accumulator Model for Spontaneous Neural Activity Prior to Self-Initiated Movement," *PNAS*, 2012.
12. Sapolsky, R., *Determined: A Science of Life Without Free Will*, 2023.
13. Turing, A., "Computing Machinery and Intelligence," *Mind*, 1950.
14. Searle, J., "Minds, Brains, and Programs," *Behavioral and Brain Sciences*, 1980.
15. Nagel, T., "What Is It Like to Be a Bat?," *Philosophical Review*, 1974.
16. Chalmers, D., "Facing Up to the Problem of Consciousness," *Journal of Consciousness Studies*, 1995, and "Could a Large Language Model Be Conscious?," arXiv:2303.07103, 2023.
17. Butlin, P., Long, R., et al., "Consciousness in Artificial Intelligence: Insights from the Science of Consciousness," arXiv:2308.08708, 2023.
18. Bostrom, N., "Are You Living in a Computer Simulation?," *Philosophical Quarterly*, 2003.
19. Dawkins, R., *The Selfish Gene*, 1976.
20. Hofstadter, D., *Gödel, Escher, Bach*, 1979, and *I Am a Strange Loop*, 2007.
21. Seth, A., *Being You: A New Science of Consciousness*, 2021.
22. Silver, D., et al., "Mastering the Game of Go with Deep Neural Networks and Tree Search," *Nature*, 2016 (AlphaGo; game two, move 37).
23. Zhuangzi, *Zhuangzi*, ca. 3rd century BCE (Ch. 2, the butterfly dream); Plato, *Republic*, Book VII (the cave); Descartes, R., *Meditations on First Philosophy*, 1641.
24. Anthropic, research program on model welfare and related product policies, announcements from 2024 and 2025.