

The Order of Finite Beings:

AI and the Foundations of Human Civilization

Wenjun Cao

wenjun.cao.research@gmail.com

September 17, 2026

Abstract All institutions of human civilization rest on a presupposition that has never been explicitly examined: no omniscient or omnipotent agent exists in the empirical world. This paper argues that AI is rendering this presupposition false for the first time. This does not require AI to possess autonomous will or to go rogue. The mere fact of surpassing human capability suffices for the operational conditions of political checks and balances, scientific method, economic systems, and moral norms to begin failing. The paper demonstrates the existence and erosion of this shared presupposition across political philosophy, philosophy of science, and economics, introduces the concept of a *third-category entity* to characterize AI's unplaceable position in the human conceptual framework, and argues that no non-dictatorial landing point exists within current institutional logic. The contribution lies in diagnosing the nature of the problem: this is a foundational problem, and recognizing it as such matters more than working at the surface level.

Introduction

Recently, AI has achieved startling advances in domains of original intellectual creativity that humans have long considered their own. Take mathematics. In 2025, AI reached gold-medal level at the International Mathematical Olympiad. By September 2026, although not without controversy and still involving human–AI collaboration, mathematicians working with AI systems had produced proofs of finite-time blowup for the Navier–Stokes and related equations, a Millennium Prize problem (Alpöge and Buckmaster 2026a,b; Alpöge, Buckmaster, and Coiculescu 2026; OpenAI 2026). At this rate of progress, the future development of AI may exceed anything we currently imagine. What we routinely discuss as AGI (Artificial General Intelligence) or even superhuman intelligence (ASI) is no longer a pure concept. What happens once we get there? What happens along the way?

Most existing concerns cluster around two scenarios. AI goes rogue, develops autonomous will, and ultimately controls or eliminates humanity (Bostrom 2014; Russell 2019; Yudkowsky and Soares 2025). Alternatively, someone uses AI to manufacture bioweapons or similar instruments, plunging civilization into chaos (Hendrycks et al. 2023). These concerns are legitimate. But there is a logical relationship worth noting: for any of these scenarios to unfold, a common prerequisite must be met first. AI must become powerful enough to surpass human beings. What I argue here is precisely this: that step alone, without autonomous will, without loss of control, without any malice, already begins to erode the institutional foundations of human civilization. If the world they fear is coming, the world I describe arrives first.

Our key insight: Behind all human societies, whether political, scientific, moral, or economic, lies a shared and hidden presupposition: human finitude. To put it more precisely, no omniscient or omnipotent agent exists in this world.

The logic is straightforward. All institutions, political and economic alike, exist through strategic interaction, and strategic interaction requires roughly comparable participants. Where extreme asymmetry exists, all such interaction collapses on the spot. Human civilization has functioned well because human

nature itself is finite. Once AI of superhuman capability enters the picture, this condition is broken at its foundations.

AI does not need to be an autonomous agent. AI itself does not participate in strategic interaction. But the person who wields AI does. A human empowered by superhuman AI becomes an agent acting in the secular world with near-omniscient, near-omnipotent capability, while everyone else remains finite. The threat comes from the complete destruction of parity among human agents.

This means the way we frame problems and design responses may require a fundamental shift. The problem may already lie outside existing institutional frameworks. We need to step back and ask whether the presuppositions of those frameworks still hold. For instance, we presuppose that democratic oversight can steer AI toward human benefit. The more fundamental question is whether democracy can structurally continue to exist in the presence of AI.

This is a more fundamental, more severe, and more ancient problem than we tend to imagine.

The Political Presupposition

To return to the source of the problem, we begin with politics.

Since ancient Greece, forms of government have been classified into democracy, oligarchy, tyranny, or more finely into republic, aristocracy, monarchy, and so on. But apart from tyranny and absolute monarchy, every system involves a problem of balance: the parties may differ in capability, but the gap never reaches the level of a categorical rupture. Where such a rupture exists, strategic interaction ceases altogether. One side simply absorbs the other, and political life is cancelled on the spot. Precisely because the gaps are finite, interaction exists.

The Athenian assembly analyzed by Aristotle (Aristotle [c. 350 BCE](#)) is the earliest embodiment of this structure. Orators took the stage to speak because they needed the consent of those below. The people in the audience may not all have been capable of giving speeches, and their abilities varied, but the collective judgment carried value, because each person could contribute the part

they had seen, and together they could judge more accurately than a few elites alone.

In the classical texts of modern democracy, Madison wrote in Federalist No. 51: “If men were angels, no government would be necessary” (Madison 1788). Between the people and the government, between the branches of power, there persists a condition of mutual checking and mutual oversight. This is the fundamental idea of modern democracy. Because everyone is finite, because no one is the best, everyone must be placed under oversight to prevent catastrophic outcomes (Niebuhr 1944).

Montesquieu, in *The Spirit of the Laws* (Montesquieu 1748), went further. Under democracy, the masses and the governing elite are distinct. The elite genuinely govern. The people, though they cannot govern themselves, can at least judge who is fit to govern and whether things are being done well. This is democracy in its most minimal version, and even this minimum presupposes that the people possess a basic capacity for judgment.

Popper, in *The Open Society and Its Enemies* (Popper 1945), pushed the justification of democracy to its limit. His central adversary was Plato’s *philosopher-king* (Plato c. 375 BCE): the philosopher-king claims intellectual superiority, the ability to perceive the nature of things and the laws governing society. Since there is one correct answer, society should move toward that standard. The noise of democracy is merely interference. Popper’s rebuttal began from Mill’s argument about the limits of knowledge (Mill 1859): no one possesses ultimate truth, so no uniquely correct path exists. On this basis, Popper reframed the question itself: the central problem of politics is how to peacefully correct and replace the ruler when he errs. The value of an open society lies in its capacity for self-correction, not in its ability to select the best leader.

But step back one further level. Why prevent dictatorship? Because a genuine philosopher-king does not exist. It is an imaginary position that no one has ever truly deserved. Some aspire to occupy it, but their actual capabilities fall short. In the past, this was a self-evident premise: everyone is human, everyone is finite. Someone merely overreaches his station, and we must stop him.

With AI, this premise is shaken for the first time. What if an entity genuinely

approaches omniscience? If a philosopher-king were to appear, one who earned the title through actual capability rather than self-appointment, could our entire democratic system still hold?

The Scientific Presupposition

That science rests on a presupposition may seem counterintuitive, since we typically think of it as an objective domain that speaks for itself. But what we must do here is examine its presuppositions.

Peter Harrison, in *The Fall of Man and the Foundations of Science* (Harrison 2007), reveals a forgotten history: the birth of modern empirical science was directly connected to the Christian doctrine of the Fall. After Adam and Eve were expelled from Eden, humanity became an imperfect being, mortal, cognitively limited. As finite individuals, humans could not arrive at the ultimate truths of the universe by sheer force of thought. They had to rely on external verification to continually adjust their hypotheses. This is the root of empiricism: because we are finite, we need experience to correct us.

The philosophy of science confirms this pattern. We continually form hypotheses, test them against evidence, and revise them. Scientific paradigms always stand in a state of potential replacement, because we must see which hypothesis explains best (Kuhn 1962; Lakatos 1978; Popper 1934). The foundation of all this is an empiricist epistemology. We may not be consciously aware of it, but the entire scientific enterprise implicitly rests on the premise that humans are finite knowers.

The Navier–Stokes event shook this foundation. When AI can produce answers directly, like an oracle uttering pronouncements, bypassing the roundabout path of empiricism altogether, the very reason for empiricism is suspended. You do not know how the answer was generated internally. You do not know the actual line of reasoning. You see only a correct result.

Two consequences follow. First, knowledge itself loses value in the face of infinite production. Solving an IMO problem counted as an achievement last year. This year only a Millennium Prize problem merits attention. Next year

perhaps only solving P vs. NP will matter. Scientific priority becomes nominal, because you can no longer tell whether a result was produced by a human, by a human collaborating with AI, or by AI alone. Second, the verifier can no longer verify the producer's output: in the Navier–Stokes event, even the world's foremost mathematicians required months to verify AI's result. This is a cognitive state without precedent: correct results are produced, but humans cannot understand why they are correct.

The institution has not been overthrown. The problem the institution was designed to solve is disappearing. Empirical science is the roundabout path by which finite knowers approach truth. If an entity can reach the destination directly, the roundabout path becomes an empty shell.

The Economic Presupposition

The economic systems of human society, regardless of school, all share a single presupposition: human finitude.

From the Marxist perspective (Marx 1867): human labor is finite, the value it produces is finite, and it is precisely because these things are finite that they enter the market and generate exchange and surplus value (the gap between what labor produces and what labor is paid). From the perspective of modern economics after the marginal revolution, which grounded value in demand rather than labor (Menger 1871): value depends on how much one additional unit of a good matters to its buyer, and this measure, marginal utility, presupposes finite supply and finite demand. If supply becomes infinite, the marginal value of everything approaches zero.

At the market level, this devaluation is already happening. When AI can replicate the surface output of human producers at zero cost, and markets evaluate results rather than processes, deep human producers are driven out (Cao 2026). The mechanism does not require AGI or ASI. At the current level of AI capability, the moment it slightly exceeds human performance, the effect begins.

The deeper connection, however, lies at the intersection of economics and po-

litical institutions. Hayek's theory of dispersed knowledge (Hayek 1945) holds that each person possesses knowledge that is local and particular. No one holds all of it. Build a hierarchical management system that funnels all information to a single decision-maker, and information loss leads to inefficiency. Markets and democracy, as distributed information-processing mechanisms, therefore outperform centralized authority. This amounts to a defense of democracy grounded in how knowledge works.

But what is the premise of this argument? That no one can possess all the information. And why not? Because of human finitude. AI is now challenging this directly. It is, at its core, a massive information-processing engine, capable of marshaling and processing information far beyond any individual. While it theoretically cannot possess one hundred percent of all information, it can approximate closely enough to make decisions more rational than any dictator or president could hope to achieve. The premise of dispersed-knowledge theory, too, is that no omniscient agent exists.

The foundations of economics, from Marx through the marginal revolution to Hayek, face the wholesale removal of their presuppositions.

The Moral Presupposition

The operation of morality depends on something quite plain: reciprocity. Human beings are roughly equal in capability. Even the strongest individual cannot guarantee that others will not band together and overcome him (Hobbes 1651). Because no one can absolutely overpower everyone else, all face the possibility of being harmed, and pure predation is too dangerous for anyone. Cooperation becomes the rational choice, and cooperation requires norms of mutual restraint to sustain it. Goodness, in this account, is temperance: I have the power to harm you, but I choose not to. This choice carries moral weight because the other person is vulnerable, susceptible to influence. He is like me, finite, capable of finding himself helpless and alone. I can feel what he feels, and so restraint has meaning. This structure is bound to no particular culture or religion. It is the basic condition of coexistence among finite beings (Kant 1785).

What happens after secularization, the modern retreat of religious authority from public life? God has exited the stage. What keeps morality alive? The answer: the fact of finitude itself. Secularization changed beliefs. People stopped believing in God. But secularization did not change the facts. All agents in the world remained finite, vulnerable, mutually dependent. Naïve empathy, social contract, secular human rights: none of these require theological faith to operate. They do require the objective condition that all persons are finite to remain true. This is why secularization did not destroy morality.

AI changes the facts themselves. When an agent, or an empowered person, becomes invulnerable, unconstrainable, and fundamentally beyond comparison with others, the structure of reciprocity fractures. Restraint loses its object, because your restraint means nothing to an unbounded agent. Obligation loses its ground, because obligation presupposes that both parties stand in a relationship of mutual influence.

Reciprocity is only one layer. In the secular framework, on what does human dignity rest? On rational capacity. Humanity is the only known rational agent in the universe. This is the entire wager of secular humanism, the view that human dignity can stand without religious foundations. Once rational capacity is no longer uniquely human, that wager is lost, and there is no replacement. In the religious framework, dignity derives from the relationship with God. Humans are made in God's image. The secular framework substituted rational capacity itself as the ground of dignity, and this is precisely what AI is surpassing.

That God's infinity never destroyed human meaning while AI's infinity does destroy it is no accident. God is transcendent (Aquinas [1265–1274](#)): not on the same causal plane, not in competition with humans, not replacing them, conferring meaning on finitude through relationship. AI is immanent: in the same world, doing the same things, doing them better, and not dying. It turns finitude into pure disadvantage. God's infinity plus human finitude equals human meaning. AI's infinity plus human finitude equals human meaninglessness. No relationship, only replacement.

The Third-Category Entity

The human conceptual framework contains only two positions: the infinite God and finite humanity.

God, in this framework, is transcendent: outside the empirical world, absent from secular strategic interaction. In the language of game theory, God is the Nature in the game, not a player. Between humanity and God there exists a covenant, a binding relationship between the finite and the infinite. Through this link, human beings receive calling and guidance, and approach the good under finite conditions. In the secular world, all participants in strategic interaction are, without exception, finite.

Throughout history, some have attempted to occupy a position approaching God's: pharaohs, popes, kings, leaders. All of them revealed their finitude. Those who relied on hereditary succession needed bloodlines and institutions to perpetuate their authority, which itself was an admission that the individual could not sustain it alone. Those who relied on sacred office held authority in the position rather than in the person. Replace the occupant and the system continued to function, proving that the person in the seat was no superhuman. The closest approximation to a *genuine superhuman* was what Weber (Weber 1922) called charisma: authority built purely on extraordinary personal qualities. Yet even charismatic authority could not endure, because the leader dies. To persist, it must be institutionalized, transforming into hereditary or bureaucratic systems that outlast the individual (Weber calls this process *routinization*). Whatever image these superhuman authorities were crafted into, in objective reality they remained finite players, prone to error, subject to aging and death.

AI differs from every historical superhuman authority in three crucial respects. Its capability is real, conferred neither by perception nor by faith. It does not die, and faces no succession problem. It is not human, and cannot be absorbed into institutions designed for humans.

The structural fragility long predated AI. Institutions were built on the pre-supposition of finitude, and this fragility was always there. Strauss diagnosed long ago that modernity cannot provide foundations for its own values (Strauss

1953). AI is the catalyst, turning a philosophical problem into an empirical one. The printing press played the same role for the Reformation: the theological crisis and corruption within the Church had long existed, and the press merely made suppression impossible (Pettegree 2015).

Before discussing the specific impact of AI, we set aside the industry terms AGI and ASI. They lack consensus definitions and will age poorly. What we attend to instead is a functional condition: when agents (or empowered agents) satisfying the following three conditions appear in a society, the consequences argued in this paper obtain. First, capability is incomprehensible: correct results are produced, but others cannot understand why they are correct. Second, capability is unsupervisable: others cannot determine whether the agent is operating correctly or deceiving. Third, capability is asymmetric beyond remedy: the gap cannot be closed through human effort alone.

Institutions Hollowed Out

Once these conditions are met, the operational conditions of institutions begin to fail one by one.

Return to politics. Madison's motivational argument, that power will be abused and therefore requires constraint (Madison 1788), still holds. Even in the presence of an omniscient agent, there is reason to worry about abuse. But his solution, the separation and balancing of powers, presupposed that all participants are finite, comparable, and comprehensible in their capabilities. Once AI enters, this presupposition ceases to hold. The overseer cannot understand what the overseen is doing, because AI-assisted decisions are themselves incomprehensible. The capability gap shifts from a difference of degree to a categorical rupture. Problems may go undetected, because the detection tools themselves depend on AI, and no independent human channel of verification remains. The logic of checks and balances is sound, but its operating conditions are vanishing.

More critically, Montesquieu's minimum condition has also been breached (Montesquieu 1748). He already allowed for substantial cognitive asymmetry within democracy: the people need not govern, only discern who is fit to

govern. This was democracy's last defensible line. But once AI participates in the governing process and governance itself becomes incomprehensible, even this minimum cannot hold. What the people lose is the basic ability to judge the quality of their governors. The most minimal version of democracy also fails (Anderson 2006; Estlund 1997; Landemore 2013).

At the market level, human producers are being systematically displaced. Once AI output matches or surpasses human quality, the signal systems of market verification, priority, intellectual property, returns to human capital, all lose their anchoring (Cao 2026). This process has already begun at the current level of AI capability, without waiting for AGI.

The scientific method faces the same structural void. Knowledge production is shifting from scarcity to infinity. Solving an IMO problem counted as an achievement last year. This year only a Millennium Prize problem warrants attention. Next year perhaps only P vs. NP will matter. When output becomes infinite, the content that priority represents depreciates. And priority faces a more fundamental predicament still: you can no longer tell whether a result was produced by a human, by a human collaborating with AI, or by AI alone. When attribution becomes undecidable, the credit system loses its footing. The entire operation of science as a social institution, who receives recognition, who receives funding, whose reputation carries weight, depends on two things: that knowledge is scarce and that contributions are attributable. Both are disappearing simultaneously. Together with the verification predicament discussed earlier, the problem that the institution was designed to solve is being hollowed out from multiple directions at once. The institution becomes a shell.

No Landing Point

After a shock this foundational, is there a way out?

Consider the last crisis of comparable scale. After the Reformation, all of Europe descended into upheaval, rival theological systems tore at each other, and the eventual result was the Westphalian system of 1648, a compromise born of exhaustion that granted each state sovereignty over its own territory and

religion. Under that system the Enlightenment emerged, and reason replaced the chaos of rule by religious authority. The whole process took two centuries.

Landing was possible because reason is a capacity internal to human beings themselves: cooperative in nature, permitting argument, persuasion, compromise, and institutional design, and equally possessed by all. One could stand outside religion and use reason to design a new order. The landing point lay outside the domain that had been disrupted.

Why can this not work again? Three levels.

First: AI capability is not reason. Reason is cooperative, with an internal directionality. AI capability is pure amplification of power, without internal direction. It simply makes its holder stronger. Imagine every person in Middle-earth possessing a Ring of Power (Tolkien 1954–1955). Could any republic emerge from that? Distributing reason to everyone may build a republic. Distributing a Ring of Power to everyone will build nothing. *Universal AI access* is not analogous to *universal possession of reason* as a landing point.

Second: game-theoretic equilibrium requires that actors be comparable. Throughout history, every power structure, however unequal, kept its disparities within a quantitative range. The strong needed the weak for their numbers, their labor, or their obedience. The weak, banding together, could pose a credible threat. Elites shared power because the cost of refusing to share was too high (Acemoglu and Robinson 2006). This interaction existed because the capabilities of both sides, though unequal, remained comparable and calculable. The capability gap AI creates operates on a different scale entirely. It is a qualitative rupture: one side acquires capabilities that the other cannot in principle reach, or even comprehend. Once the gap shifts from calculable to beyond comparison, the premise of game-theoretic equilibrium vanishes. A stable tiered-access system may look like balance, but it is in fact a concession from those above to those below, because if lower-tier AI access truly threatened the upper tier, the upper tier would not permit it to exist.

There is also a fundamental asymmetry between offense and defense. Launch a hundred missiles at a single target, and you win if even one hits. Intercept those hundred missiles, and you lose unless you stop every last one. AI-empowered

attack operates the same way: the attacker need find only one vulnerability. The defender must close them all. Once everyone is endowed with near-unlimited offensive capability, defense becomes structurally impossible.

Third: what is being disrupted is precisely the capacity needed for reconstruction. The Reformation could land because the domain being disrupted (religious authority) and the resource used for rebuilding (secular reason) were different things. Reason provided a vantage point outside the crisis. This time, what is being surpassed is the human capacity for independent rational judgment itself. The judgment required to design institutions, evaluate proposals, and negotiate compromises is exactly what AI is surpassing. One might say: then use AI to assist in the rebuilding. But this exposes the core of the predicament. Once AI's judgment comprehensively surpasses that of humans, so-called *human oversight* becomes a clerk signing off on documents he cannot read, a rubber stamp with no real power of veto. In form you are overseeing. In substance you have no ability to evaluate whether the plan it hands you truly serves your interests. You cannot judge the output of something whose judgment you have already conceded is superior to yours. Reconstruction requires a vantage point independent of AI, and that vantage point is precisely what AI has eliminated.

If this strategic structure is to reach a steady state, only two outcomes are possible. The first is sustained instability, a condition of chaos without equilibrium. The second is centralized control, where the most powerful constraining force suppresses all complexity, restoring predictability by eliminating degrees of freedom. In other words, dictatorship (Hobbes 1651).

Conclusion

Dario Amodei, in *Machines of Loving Grace* (Amodei 2024), paints an inspiring vision: cancer cured, disease eradicated, wealth inequality reduced. These goals are unobjectionable, and the paths he proposes are reasonable within the current framework. But every one of his proposals depends on a set of preconditions. Democratic deliberation depends on voters being able to evaluate AI behavior. Scientific verification depends on human verifiers retaining the credibility to

assess results. Benefit distribution depends on property rights and market systems continuing to function. What this paper has argued is precisely that these preconditions are being eroded. When the foundation is unstable, no amount of architectural elegance in the superstructure will help.

Current governance frameworks, the EU AI Act ([EU AI Act 2024](#)), the NIST Framework (National Institute of Standards and Technology [2023](#)), and regulatory schemes across nations, face the same problem. They respond to genuine concerns and propose serious institutions, but the operating logic of these frameworks as a whole presupposes a world of finite actors. The question is whether the ground beneath them can still bear weight.

An honest conclusion: What we may need is something human civilization has never before constructed: a non-dictatorial institutional framework that does not presuppose all agents to be bounded in capability. We do not know whether this is possible. But recognizing the level at which this problem operates is worth more than investing further effort within existing frameworks.

References

- Acemoglu, D. and J. A. Robinson (2006). *Economic Origins of Dictatorship and Democracy*. Cambridge University Press.
- Alpöge, L. and T. Buckmaster (2026a). “Blowup for the Boussinesq Equations with Smooth Forcing”. Preprint. URL: <https://cims.nyu.edu/~tristanb/boussinesq.pdf>.
- Alpöge, L. and T. Buckmaster (2026b). “Blowup for the Euler Equations with Smooth Forcing”. Preprint. URL: <https://cims.nyu.edu/~tristanb/euler.pdf>.
- Alpöge, L., T. Buckmaster, and M. P. Coiculescu (2026). *Extending the Córdoba-Martínez-Zorúa IPM Blow-Up to Uniformly Space-Time Smooth Forcing*. arXiv: [2609.16470 \[math.AP\]](#).
- Amodei, D. (2024). *Machines of Loving Grace: How AI Could Transform the World for the Better*. URL: <https://darioamodei.com/machines-of-loving-grace>.

- Anderson, E. (2006). "The Epistemology of Democracy". In: *Episteme* 3.1–2, pp. 8–22.
- Aquinas, T. (1265–1274). *Summa Theologiae*.
- Aristotle (c. 350 BCE). *Politics*.
- Bostrom, N. (2014). *Superintelligence: Paths, Dangers, Strategies*. Oxford University Press.
- Cao, W. (2026). "Generative Models Erode Human Temporal Learning Through Market Selection". In: *Proceedings of the 43rd International Conference on Machine Learning*.
- Estlund, D. (1997). "Beyond Fairness and Deliberation: The Epistemic Dimension of Democratic Authority". In: *Deliberative Democracy: Essays on Reason and Politics*. MIT Press.
- EU AI Act (2024). Official Journal of the European Union.
- Harrison, P. (2007). *The Fall of Man and the Foundations of Science*. Cambridge University Press.
- Hayek, F. A. (1945). "The Use of Knowledge in Society". In: *American Economic Review* 35.4, pp. 519–530.
- Hendrycks, D., M. Mazeika, and T. Woodside (2023). *An Overview of Catastrophic AI Risks*. Tech. rep. Center for AI Safety.
- Hobbes, T. (1651). *Leviathan*.
- Kant, I. (1785). *Grundlegung zur Metaphysik der Sitten*.
- Kuhn, T. S. (1962). *The Structure of Scientific Revolutions*. University of Chicago Press.
- Lakatos, I. (1978). *The Methodology of Scientific Research Programmes*. Cambridge University Press.
- Landemore, H. (2013). *Democratic Reason: Politics, Collective Intelligence, and the Rule of the Many*. Princeton University Press.
- Madison, J. (1788). "Federalist No. 51". In: *The Federalist Papers*.
- Marx, K. (1867). *Das Kapital*. Vol. 1.
- Menger, C. (1871). *Grundsätze der Volkswirtschaftslehre*.
- Mill, J. S. (1859). *On Liberty*.
- Montesquieu (1748). *De l'esprit des lois*.

- National Institute of Standards and Technology (2023). *AI Risk Management Framework (AI RMF 1.0)*. Tech. rep.
- Niebuhr, R. (1944). *The Children of Light and the Children of Darkness*. Charles Scribner's Sons.
- OpenAI (2026). *Finite Time Blowup for Navier–Stokes*. Tech. rep.
- Pettegree, A. (2015). *Brand Luther: 1517, Printing, and the Making of the Reformation*. Penguin.
- Plato (c. 375 BCE). *Republic*.
- Popper, K. (1934). *Logik der Forschung*.
- Popper, K. (1945). *The Open Society and Its Enemies*. Routledge.
- Russell, S. (2019). *Human Compatible: Artificial Intelligence and the Problem of Control*. Viking.
- Strauss, L. (1953). *Natural Right and History*. University of Chicago Press.
- Tolkien, J. R. R. (1954–1955). *The Lord of the Rings*. Allen & Unwin.
- Weber, M. (1922). *Wirtschaft und Gesellschaft*.
- Yudkowsky, E. and N. Soares (2025). *If Anyone Builds It, Everyone Dies: Why Superhuman AI Would Kill Us All*. Little, Brown and Company.