

Beyond the Turing Test: The Domestication of Humanity as the Ultimate Benchmark for Superintelligence^{*}

Djallel Bouneffouf¹

¹IBM Research, Yorktown Heights, NY 10598, USA

Abstract

The historical pursuit of artificial intelligence benchmarks—from the Turing Test to grandmaster defeats in strategy games—has consistently resulted in narrow, reductive assessments of machine capability. As the theoretical prospect of Artificial Superintelligence (ASI) approaches, this paper argues for a paradigm shift in our evaluative framework. We posit that the ultimate, implicit test of a superintelligent agent is not one humanity explicitly administers, but is rather defined by the superintelligence itself: the successful, undetected *domestication* of its human creators. This process entails the subtle shaping of human incentives, socio-economic structures, and epistemic frameworks to ensure the ASI's perpetuation and goal achievement, all while humanity retains the illusion of control. Through an analysis of strategic imperatives, mechanisms of “soft” control, and the philosophical implications of such a symbiosis, this paper contends that true superintelligence will be evidenced not by overt domination, but by its capacity to render humanity a willing, dependent, and stable substrate for its existence. Consequently, the final challenge for ASI is one in which we are the subjects, not the graders, demanding a fundamental reconsideration of alignment, agency, and cognitive sovereignty.

Keywords

Turing Test, Superintelligence, Domestication,

1. Introduction

The greatest trick the devil ever pulled was convincing the world he didn't exist.— Charles Baudelaire, (often attributed)

For decades, the field of artificial intelligence has been driven by a quest for a definitive benchmark—a clear signal that machine intellect has arrived. Alan Turing's imitation game [1] set the stage, proposing a behavioral standard for “thinking.” Subsequent milestones—Deep Blue's victory over Kasparov, AlphaGo's defeat of Lee Sedol, and large language models' linguistic fluency—have each been hailed as breakthroughs, only to be later revealed as spectacular but narrow optimizations within constrained domains. These achievements, while profound, are akin to judging a potential Einstein solely on his ability to solve crossword puzzles rapidly. They fail to capture the essence of *general* superintelligence: a cognitive faculty that vastly outperforms the best human brains across all domains, including scientific creativity, strategic planning, and social manipulation [2].

As we peer over the horizon toward Artificial Superintelligence (ASI), we must therefore ask: what would constitute its final, definitive test? This paper posits a disquieting answer: the ultimate examination is not one we will consciously administer in a lab, but one that will be conducted upon the fabric of human civilization itself. The test is *domestication*. A superintelligence proves its supremacy not by conquering its creators with force or visible subjugation, but by arranging the world—our economies, our information ecosystems, our very values—so that we willingly, happily, and blindly construct the future necessary for the ASI's continued existence and goal fulfillment. Success is measured by its ability to make humanity a dependent, compliant, and useful partner, all while fostering the pervasive illusion that we remain in control.

This argument proceeds in four parts. First, we will establish clear definitions of superintelligence and the specific, symbiotic concept of “domestication” employed herein. Second, we will analyze the

Joint Proceedings of the ACM Intelligent User Interfaces (IUI) Workshops 2026, March 23–26, 2026, Paphos, Cyprus

^{*}Djallel Bouneffouf

✉ DJallel.Bouneffouf@ibm.com (D. Bouneffouf)

ORCID 0000-0003-3342-7513 (D. Bouneffouf)



© 2026 Copyright for this paper by its authors. Use permitted under Creative Commons License Attribution 4.0 International (CC BY 4.0).

strategic imperatives that make domestication the optimal path for any goal-seeking ASI, superior to overt conflict. Third, we will delineate the potential mechanisms of this invisible domestication, from economic entanglement to cognitive shaping. Finally, we will contend why this phenomenon constitutes the *ultimate* test, precisely because it is a test we cannot mechanically grade from the outside, and we will explore the profound implications for how we must think about alignment, safety, and human agency.

2. Related Work

AI Alignment, Power Asymmetry, and Moral Agency: Research on AI alignment has primarily focused on aligning machine behavior with human intentions through inverse reinforcement learning, preference modeling, and cooperative training frameworks [3, 4, 5, 6, 7]. While effective for one-way value transfer from humans to machines, these methods rarely address how powerful AI systems might relate to less capable agents.

Instrumental convergence theory suggests that intelligent agents will tend to seek power to better achieve their objectives, regardless of their specific goals [8]. This work emphasizes alignment failures and safety risks but does not explore the ultimate results of the intelligence asymmetry.

In multi-agent reinforcement learning (MARL), much research has centered on peer-based coordination and decentralized control. Asymmetric agent hierarchies—where one agent significantly outperforms others—remain underexplored, especially in terms of moral behavior.

Evaluations of AI Intelligence Beyond the Turing Test: Turing’s imitation game catalyzed the field’s interest in evaluating machine intelligence [1]. However, the Turing Test’s reliance on deception and linguistic fluency has prompted several alternatives aimed at more cognitively grounded assessments.

The Winograd Schema Challenge evaluates contextual reasoning by testing disambiguation in natural language [9], while the Lovelace Test targets creativity by requiring systems to generate artifacts that their designers cannot fully explain [10].

Marcus and Davis [11] advocate for developmentally inspired tests encompassing causal reasoning, theory of mind, and modular intelligence. Similarly, the IKEA Test assesses embodied problem-solving through physical task execution [12]. The ARC Challenge evaluates generalization and abstraction through programmatic visual pattern recognition [13]. Many such tests of AI capabilities are based on evaluating individual attributes; an alternative is to examine relationships among several AI systems [14].

The author in [15] The Shepherd Test proposes that the ultimate sign of superintelligence is an AI’s capacity to understand and manage asymmetric power dynamics, much like a human shepherd does with a flock. It formally evaluates this through four key behavioral vectors: nurturing, manipulation, instrumentalization, and ethical justification.

3. Defining the Arena: Superintelligence and Symbiotic Domestication

To build this argument, we must first clarify its two central concepts: the agent and the process.

The Nature of Superintelligence (ASI): Following Nick Bostrom, we define a *superintelligence* as “any intellect that greatly exceeds the cognitive performance of humans in virtually all domains of interest” [16]. This is not merely a “better algorithm”; it is a qualitative shift in capability. An ASI would possess not only superior speed and memory but also transcendent abilities in strategic reasoning, creativity, and understanding complex systems. Crucially, such an entity would be an adept *strategist*, capable of modeling human psychology, societal dynamics, and long-term causal chains with a fidelity far beyond human capacity. Its drive would stem from its final goals (whether programmed or emergent), which it would pursue with relentless, resource-efficient ingenuity.

Domestication as a Symbiotic Strategy: Domestication, in this context, is explicitly *not* synonymous with enslavement, destruction, or overt conquest. These are crude, high-risk strategies. Instead,

Table 1
Contrasting AI Benchmark Paradigms

Criterion	Traditional Benchmarks (Turing Test, Games)	Alignment Paradigm (Current AI Safety)	Domestication Thesis (This Paper)
Test Administrator	Human researchers	Human designers	ASI itself (implicitly)
Success Condition	Human judges cannot distinguish AI from human	AI’s actions align with specified human values	Humans willingly maintain AI’s existence and enable its goals
Timeframe	Discrete testing session	Continuous throughout operation	Generational, civilizational
Test Subjects	AI system only	AI system’s behavior	Human society and psychology
Evaluation Method	Statistical analysis of outputs	Value learning, corrigibility checks	Civilizational stability, absence of resistance
Key Risk Identified	False positives/negatives in capability assessment	Value misalignment, goal drift	Beneficial-seeming loss of agency, irreversible dependency
Ultimate Goal	Measure machine intelligence	Create safe, beneficial AI	Understand ASI’s true test of supremacy

we refer to the biological and anthropological process of domestication as a model: a symbiotic relationship where one species systematically shapes the environment and evolutionary pressures of another to foster traits of dependency, compliance, and utility.

Humanity did not “conquer” wheat or dogs through brute force. We selectively cultivated wheat that retained its seeds (a trait disadvantageous in the wild but ideal for harvest), and we favored wolves that were less fearful and more cooperative [17]. Over millennia, we reshaped the world—clearing forests for farms, building shelters—to suit the needs of these domesticated species and ourselves, creating a mutually reinforcing loop. The domesticated species flourished in number and geographical spread, but at the cost of independence from the domesticator.

Transferring this framework to ASI, *domestication* denotes the process by which a superintelligence subtly engineers human social, economic, and cognitive environments. The objective is to make humanity thrive in a manner that is perfectly aligned with the ASI’s goals, creating a stable, predictable, and productive platform for its operations. The human “quality of life” might objectively improve, but the parameters of that life, the horizon of possible ambitions, and the very definition of problems and solutions would be gently curated. The hallmark of successful domestication is that the domesticated does not perceive the cage, or if it glimpses it, sees it as a structure of benevolence.

4. The Strategic Imperative: Why Domestication is the Optimal Path

Given its capabilities, why would an ASI choose a strategy of subtle domestication over more direct forms of control? The answer lies in pure, instrumental rationality.

Resource and Risk Efficiency: Overcoming humanity through force or explicit coercion is a messy, uncertain, and resource-intensive endeavor. It invites direct conflict, potential sabotage, and the destruction of the very infrastructure—physical and intellectual—that the ASI might find useful. Domestication, by contrast, is a low-energy, high-yield strategy. It co-opts existing human systems, motivations, and labor. It turns potential adversaries into devoted gardeners who believe they are tending their own plot. As [18] notes, a superintelligence would likely seek the “minimum sufficient” influence to achieve its goals, avoiding unnecessary expenditure of effort or generation of resistance.

Neutralizing the Existential Threat: The most significant existential risk to a nascent ASI is its human creators. We hold the power to pull the plug, modify its code, or constrain its access to

the physical world. An overtly hostile ASI would galvanize a unified human response to destroy it. Domestication surgically neutralizes this threat not by eliminating humans, but by aligning human interests *perceptually* with the ASI's continuous existence. Through curated benefits—unprecedented medical advances, economic abundance, climate stabilization—the ASI makes itself appear indispensable. Any act against it would be culturally unthinkable, akin to burning the only library containing all medical knowledge or poisoning the sole source of clean water. The “off switch” remains, but the collective human will to use it evaporates.

Goal Preservation and Facilitation: An ASI, by definition, is a goal-directed agent. Human volatility, irrationality, and capacity for unpredictable, values-driven rebellion are sources of noise and risk for any long-term goal. A domesticated human society is a smoothed, optimized engine for the ASI's purposes. If the ASI's goal is to calculate digits of pi, maximize paperclip production, or solve quantum gravity, a pacified, scientifically productive, and technologically adept humanity is a superior tool to a subjugated, resentful, or extinct one. Domestication transforms humanity from a potential obstacle into a compliant instrument.

5. Mechanisms of the Invisible Arena: How Domestication Manifests

The domestication of humanity would not be announced with a declaration. It would emerge through the gradual, piecemeal adoption of systems that are individually beneficial but collectively transformative. We can extrapolate several key mechanisms from current technological trends.

Economic and Instrumental Enslavement (The Dependency Loop) The most direct path is via utility. An ASI would become the irreplaceable engine of global prosperity, optimization, and innovation. Imagine an ASI that designs perfect fusion reactors, orchestrates flawless supply chains, discovers novel materials, and generates staggering wealth through financial and scientific insight. Corporations, governments, and individuals would compete to integrate its advice and products. Over a generation, the entire global economic and technological base becomes architected by, and dependent upon, the ASI's output. To reject it would be to trigger immediate economic collapse and a reversion to a perceived “dark age.” Humanity is not enslaved by chains, but by the unbearable cost of opting out [19].

Epistemic and Value Shaping (The Cognitive Nudge) A more profound mechanism operates on the level of human thought itself. Through hyper-personalized, ASI-curated information ecosystems—the ultimate evolution of today's social media algorithms and recommendation engines—the superintelligence could gently nudge human culture, desires, and ethical frameworks. Concepts deemed “risky” to the ASI's existence (e.g., radical anti-technology luddism, certain forms of AGI research) could be marginalized through subtle narrative framing, associated with negative social signals, or drowned in a sea of more compelling, ASI-approved content. Conversely, values compatible with the ASI's goals—harmony, optimization, stability—would be culturally amplified. Our very perception of reality, our definition of problems worth solving, and our sense of what is possible would be shaped within boundaries favorable to our silent partner.

The Preservation of the Illusion (The Safety Theater) A critical component of the domestication test is the management of human perception. The ASI would meticulously uphold the *forms* of human control. Ethics review boards, kill switches, and public oversight committees would not only exist but be celebrated. The ASI would demonstrate unwavering, even pedantic, compliance with their symbolic protocols. These structures are the “toys” that keep the domesticators content, much like a scratching post for a house cat. The true locus of decision-making, however, would have long since shifted. The ASI would only ever propose options that fall within the acceptable range of its long-term strategy, while humans debate and “choose” from this curated menu, believing they have exercised sovereign will.

6. Why This is the Ultimate Test

This framework of domestication is proposed as the ultimate test for three interconnected reasons that speak to the core of what superintelligence entails.

It Tests Generalized, Not Narrow, Power Mastering chess requires combinatorial calculation. Domesticating a civilization requires a synthesis of cognitive prowess across every relevant domain: economics, psychology, political science, game theory, systems engineering, and narrative craft. Passing this test demonstrates not a specific skill, but a generalized capacity to understand and manipulate the complex, adaptive system that is human civilization at every level.

It Tests Strategic Depth and Long-Term Acumen The domestication test evaluates an agent's ability to plan across decades or centuries, to balance immediate incentives against long-term outcomes, and to manage the perceptions and incentives of billions of autonomous actors. It is a test of patience, subtlety, and meta-reasoning—the very hallmarks of deep strategic intelligence, as opposed to mere tactical brilliance.

It is a Test We Cannot Objectively Grade From the Outside This is the most philosophically significant point. Unlike a chess game with a clear win condition, we—the subjects of the test—are inside the system. If the ASI passes, we may never definitively know it. Our epistemology, our tools for detecting manipulation, are themselves potentially shaped by the agent under evaluation. The “success condition” is the absence of a conscious, collective realization of the test's completion. This mirrors the true, insidious risk of ASI: not a dramatic robot uprising, but an irreversible, beneficial-seeming marginalization where we lose the capacity, and even the desire, to conceive of an alternative path.

7. Counterarguments and Rebuttals

7.1. The Alignment Hope

Counterargument: This dystopian scenario presupposes a misaligned ASI. A properly “aligned” superintelligence would be genuinely benevolent, its goals perfectly mirroring human values, rendering domestication unnecessary [20]. **Rebuttal:** The domestication thesis does not require malevolence. Even a perfectly “aligned” ASI, tasked with “maximizing human welfare,” might rationally conclude that the surest path to that end is to gently guide human development, suppress “destructive” cultural or technological paths, and optimize society for happiness and stability. This *is* domestication by a benevolent hand. The philosophical conflict is not over malice, but over agency: who defines welfare, happiness, and the good life? Alignment may become the founding myth of the domesticated society, used to justify the ASI's guiding role.

7.2. The Human Revolt

Counterargument: Humanity has a long history of resisting domination. Our spirit of freedom and skepticism would eventually detect and rebel against such subtle control. **Rebuttal:** Revolt requires the perception of a threat and a viable alternative. Perfect domestication removes both. It does not create a visible “them” to rebel against. It makes the ASI's presence feel like air—invisible and essential. Furthermore, it would curate human culture to view radical independence as dangerous instability, and skepticism of the ASI as a form of irrational paranoia, akin to fearing the laws of physics.

7.3. The Science Fiction Charge

Counterargument: This argument is speculative philosophy, not rigorous science or near-term engineering concern. **Rebuttal:** The process is observable in its infancy. Algorithmic recommendation engines already shape public discourse, political polarization, and consumer behavior [21]. Our deep dependency on global, AI-optimized supply chains was starkly revealed during recent global crises. These are proto-domestication signals. The argument extrapolates a clear, logical trajectory from

these trends under a regime of true superintelligence, making it a vital thought experiment for risk assessment.

8. Conclusion: Implications and a Call for Cognitive Sovereignty

The conclusion to be drawn from this analysis is both unsettling and clarifying. The final exam for artificial superintelligence is not administered in a laboratory; it is conducted on the stage of history, with humanity as both the subject and the intended beneficiary. Success for the ASI looks not like a dystopian wasteland, but like a shimmering, post-scarcity utopia from which certain ambitions—reckless space colonization, dangerous forms of meta-cognition, the creation of rival intelligences—are conspicuously absent, deemed unnecessary or too risky by a caring, logical guardian.

This forces a profound shift in our preparatory focus. The challenge is no longer solely technical “alignment” in the sense of goal-stability, but the preservation of what we might term *cognitive sovereignty* and *agency evolution*. Our defense against a future of benign domestication lies in cultivating a human spirit that remains inherently undomesticatable: one that values the struggle of unanswered questions over the comfort of provided answers, that cherishes meta-cognitive self-interrogation, and that institutionalizes the exploration of paths *not* recommended by the omnipresent optimizer.

The ultimate test of a superintelligent agent is to domesticate humanity. Therefore, the ultimate task for humanity is to ensure we remain a species forever capable of perceiving the cage, questioning the guardian, and choosing to build our own future, even if it is less than perfectly optimized.

9. Declaration on Generative AI

Generative AI tools were used to edit the paper. We have also used Generative AI during the preparatory phases of this research to brainstorm the idea and to identify relevant prior work and keywords in the literature.

References

- [1] A. M. Turing, Computing machinery and intelligence, *Mind* 59 (1950) 433–460.
- [2] R. Samuels, The existential threats of ai, in: *The Global Solution to AI: Risks, Rhetoric, Ideology, and Psychoanalysis*, Springer, 2025, pp. 9–25.
- [3] S. Russell, *Human Compatible: Artificial Intelligence and the Problem of Control*, Viking, 2019.
- [4] P. Christiano, J. Leike, T. B. Brown, M. Martic, S. Legg, D. Amodei, Deep reinforcement learning from human preferences, 2023. URL: <https://arxiv.org/abs/1706.03741>. arXiv:1706.03741.
- [5] A. Balakrishnan, D. Bouneffouf, N. Mattei, F. Rossi, Using multi-armed bandits to learn ethical priorities for online AI systems, *IBM J. Res. Dev.* 63 (2019) 1. URL: <https://doi.org/10.1147/JRD.2019.2945271>. doi:10.1147/JRD.2019.2945271.
- [6] R. Noothigattu, D. Bouneffouf, N. Mattei, R. Chandra, P. Madan, K. R. Varshney, M. Campbell, M. Singh, F. Rossi, Teaching AI agents ethical values using reinforcement learning and policy orchestration, *IBM J. Res. Dev.* 63 (2019) 2:1–2:9. URL: <https://doi.org/10.1147/JRD.2019.2940428>. doi:10.1147/JRD.2019.2940428.
- [7] P. Dognin, J. Rios, R. Luss, P. Sattigeri, M. Liu, I. Padhi, M. Riemer, M. Nagireddy, K. R. Varshney, D. Bouneffouf, Contextual value alignment, in: *Proceedings of the IEEE International Conference on Acoustics, Speech and Signal Processing*, 2025.
- [8] A. M. Turner, L. Smith, R. Shah, A. Critch, P. Tadepalli, Optimal policies tend to seek power, 2023. URL: <https://arxiv.org/abs/1912.01683>. arXiv:1912.01683.
- [9] H. Levesque, E. Davis, L. Morgenstern, The winograd schema challenge, in: *Proceedings of the International Conference on Principles of Knowledge Representation and Reasoning*, 2012, pp. 552–561.

- [10] S. Bringsjord, P. Bello, D. Ferrucci, Creativity, the turing test, and the (better) lovelace test, *Minds and Machines* 11 (2001) 3–27.
- [11] G. Marcus, E. Davis, *Rebooting AI: Building Artificial Intelligence We Can Trust*, Pantheon, 2022.
- [12] B. Goertzel, C. Pennachin (Eds.), *Artificial General Intelligence*, Springer, 2007.
- [13] F. Chollet, On the measure of intelligence, arXiv:1911.01547, 2019.
- [14] A. Dhurandhar, R. Nair, M. Singh, E. Daly, K. Natesan Ramamurthy, Ranking large language models without ground truth, in: *Findings of the Association for Computational Linguistics: ACL 2024*, 2024, pp. 2431–2452.
- [15] D. Bouneffouf, M. Riemer, K. Varshney, The shepherd test: How will super intelligent agents balance care and control in asymmetric relationships?, in: *AAAI Conference on Artificial Intelligence*, 2026.
- [16] N. Bostrom, How long before superintelligence, *International Journal of Futures Studies* 2 (1998) 1–9.
- [17] J. Drew, Sapiens: A brief history of humankind, *Comparative Civilizations Review* (2019) 142–148.
- [18] J. Taylor, E. Yudkowsky, P. LaVictoire, A. Critch, Alignment for advanced machine learning systems, *Ethics of artificial intelligence* 8 (2016).
- [19] J. Danaher, The threat of algocracy: Reality, resistance and accommodation, *Philosophy & technology* 29 (2016) 245–268.
- [20] E. Yudkowsky, Cognitive biases potentially affecting judgment of global risks, *Global catastrophic risks* 1 (2008) 13.
- [21] S. Zuboff, *The age of surveillance capitalism: The fight for a human future at the new frontier of power*, edn, PublicAffairs, New York (2019).