



Expert Insights
PERSPECTIVE ON A TIMELY POLICY ISSUE

JOEL B. PREDD, BENJAMIN BOUDREAUX, MATT CHESSEN, BEBA CIBRALIC, EDWARD GEIST,
KAMARIA HORTON, AMANDA KERRIGAN, WILLIAM MARCELLINO, SHANSHAN MEI,
JARED MONDSCHIN, ALVIN MOON, TOBIAS SYTSMA

A U.S. Strategy to Secure Geopolitical Advantage on an Uncertain Path to Superintelligence

Maintaining Freedom of Action

For more information on this publication, visit www.rand.org/t/PEA5105-1.

About RAND

RAND is a research organization that develops solutions to public policy challenges to help make communities throughout the world safer and more secure, healthier and more prosperous. RAND is nonprofit, nonpartisan, and committed to the public interest. To learn more about RAND, visit www.rand.org.

Research Integrity

Our research integrity is grounded in RAND's core values of quality and objectivity. Rigorous quality assurance procedures, conflict of interest screening, and transparency in funding ensure that every study is objective and nonpartisan. Learn more at www.rand.org/integrity.

RAND's publications do not necessarily reflect the opinions of its research clients and sponsors.

Published by the RAND Corporation, Santa Monica, Calif.

© 2026 RAND Corporation

RAND® is a registered trademark.

Limited Print and Electronic Distribution Rights

This publication and trademark(s) contained herein are protected by law. This representation of RAND intellectual property is provided for noncommercial use only. Unauthorized posting of this publication online is prohibited; linking directly to its webpage on rand.org is encouraged. Permission is required from RAND to reproduce, or reuse in another form, any of its research products for commercial purposes. For information on reprint and reuse permissions, visit www.rand.org/about/publishing/permissions.

About This Paper

Frontier artificial intelligence (AI) continues to advance, and an increasing number of experts contend that the creation of superintelligence is possible. Although progress might slow or even halt, the transformative potential of superintelligence demands proactive strategy-making. The stakes are high, not only for geopolitics, but for humanity's survival and agency. What, then, should U.S. strategy be? Should the United States pursue AI supremacy by monopolizing the frontier? Co-develop advanced AI with rivals? Deter its advance? Pause and freeze its development? And in a competitive environment, how will other powers and geopolitical rivals respond to the choices they face?

In this paper, we argue that seven archetypal strategies dominate the debate, that current evidence cannot resolve five pivotal uncertainties about the most feasible and sensible strategy, and that committing prematurely to any one of them, or deferring commitment too long, forecloses options. We therefore recommend a **Freedom of Action strategy** to build and preserve the ability to secure geopolitical advantage and ensure humanity's survival with agency through the transition to superintelligence.

The strategy calls for investments in (1) building a human-AI ecosystem that preserves human survival and agency, (2) an AI security architecture, (3) adaptations of the legacy national security enterprises for the AI era, and (4) the capacity of citizens, firms, and governments to respond. The strategy requires active monitoring for indications and warnings of changes in the technological and strategic environments. Above all, it requires monitoring for movement in the five pivotal uncertainties, guarding against catastrophe and lock-in, and seizing geopolitical advantage. Freedom of action is the guiding policy rather than the goal; the ends remain U.S. geopolitical advantage and humanity's survival with agency. Preserving the capacity to choose is, under present uncertainty, the most reliable means of securing geopolitical advantage, survival, and agency. Because much of what must be built sits with the private sector, the strategy's success depends on fostering an effective partnership between the U.S. government, AI companies, businesses, and civil society.

Center for the Geopolitics of Artificial General Intelligence

RAND Global and Emerging Risks is a division of RAND that delivers rigorous and objective public policy research on the most consequential challenges to civilization and global security. This work was undertaken by the division's Center for the Geopolitics of Artificial General Intelligence (AGI), which is committed to helping decisionmakers understand, anticipate, and prepare to navigate the national security and geopolitical implications of AGI. The center convenes leading technologists, strategists, economists, political scientists, and outside experts to consider the feasibility and effectiveness of prospective AGI-enabled capabilities; the domestic and international implications of their use; and the strategies and policies that governments, businesses, and civil society could adopt to respond to new realities. For more information, visit www.rand.org/geopolitics-of-agi.

Funding

This effort was independently initiated and conducted within the Center for the Geopolitics of Artificial General Intelligence using income from operations and gifts from RAND supporters, including philanthropic gifts made or recommended by DALHAP Investments Ltd., Ergo Impact, Founders Pledge, Charlottes och Fredriks Stiftelse, Good Ventures, Longview, and Coefficient Giving. RAND donors and grantors have no influence over research findings or recommendations.

Acknowledgments

We thank Greg McElvey, Helen Toner, Andy Hoehn, and Jim Baker for their thoughtful and constructive reviews, which substantially improved this paper. We are also grateful to our colleagues in RAND's Center for the Geopolitics of AGI and across RAND, whose ideas and feedback shaped this work throughout its development.

We used frontier large language models (LLMs) and LLM project knowledge bases to organize and synthesize several hundred pages of research memorandums produced by the project team and to draft initial text for portions of this paper. The tool was not used as a source of facts. All factual claims were verified against cited sources, and all text was drafted or substantially revised by the authors, who are responsible for the analysis, findings, and recommendations presented here.

Summary

The possible emergence of superintelligence is a strategic challenge without precedent, and the United States has no settled strategy for it. Seven archetypal strategies dominate the current debate, but the choice among them depends on questions that present evidence cannot resolve. In this paper, we make five claims (followed by one recommendation):

- **Seven archetypal strategies** are available for superintelligence, each with a coherent win state, theory of victory, and structural pathology.
- **Five pivotal uncertainties** determine which strategies may be available or favored: how close the danger is, whether humans and AI can coexist, whether cooperative restraint is possible, whether a decisive advantage is attainable, and whether rival programs can be suppressed.
- **Commitment cuts both ways:** Committing prematurely to one strategy, or deferring commitment too long, forecloses the options available.
- **Capabilities should be built before they are needed:** Switching strategies once a crisis is underway is slow, costly, and potentially destabilizing, and the capabilities that a strategy depends on take years to build, so the capacity to commit must be in place before the moment of need.
- **Investments send signals:** The choices that the United States makes shape rivals' assumptions before any strategy is officially declared, so signals should be chosen deliberately, to widen U.S. options and move rivals toward restraint.

We therefore recommend a **Freedom of Action strategy**: Build and preserve freedom of action to secure geopolitical advantage and ensure humanity's survival with agency through the transition to superintelligence. The strategy calls for investments in (1) building a human-AI ecosystem, (2) an AI security architecture, (3) adaptations of the legacy national security enterprises for the AI era, and (4) the capacity of citizens, firms, and governments to respond.

The strategy requires active monitoring for indications and warnings of changes in the technological and strategic environments, guarding against catastrophe and lock-in, and seizing advantages when they are available. The United States should commit to one of the archetypal strategies when the evidence shows that it is compelling or when a strategy's latest decision date arrives with the evidence still unresolved. Freedom of action is the guiding policy rather than the goal; the ends remain U.S. geopolitical advantage and humanity's survival with agency, and preserving the capacity to choose is, under present uncertainty, the most reliable means of securing them. Because much of what must be built sits with the private sector, the strategy's success depends on fostering an effective partnership between the U.S. government, AI companies, businesses, and civil society.

Contents

About This Paper	iii
Summary	v
Figures and Table	vii
A U.S. Strategy to Secure Geopolitical Advantage on an Uncertain Path to Superintelligence: Maintaining	
Freedom of Action	1
Diagnosis and Prescription	2
Seven Archetypal Strategies	3
Five Pivotal Uncertainties	6
Risks of Foreclosure	9
AI Strategy in a Competitive Environment	12
A Strategy to Maintain Freedom of Action	13
Resources and Trade-Offs	21
Critiques	22
Conclusion	23
Bibliography	25
About the Authors	29

Figures and Table

Figures

Figure 1. How the Five Pivots Resolve to Seven Strategies	8
Figure 2. A Freedom of Action Strategy	20

Table

Table 1. Seven Archetypal Strategies for Frontier AI	5
--	---

A U.S. Strategy to Secure Geopolitical Advantage on an Uncertain Path to Superintelligence: Maintaining Freedom of Action

Frontier artificial intelligence (AI) continues to advance, and an increasing number of experts contend that the creation of superintelligence is possible (Amodei, 2024; Kokotajlo et al., 2025).¹ Although progress could slow or even halt, the transformative potential of superintelligence demands proactive strategy-making. The stakes are high, not only for geopolitics but also for humanity's survival and agency.

What, then, should U.S. strategy be? Should the United States pursue AI supremacy by monopolizing the frontier? Co-develop advanced AI with rivals? Deter its advance? Pause and freeze its development? And in a competitive environment, how will other powers and geopolitical rivals respond to the same choice?

In this paper, we make five claims:

- First, seven archetypal strategies are available for governing frontier AI as it advances along an uncertain path to superintelligence, each embodying different win states and characteristic modes of failure.
- Second, the sensibility and feasibility of each strategy hinge on five pivotal uncertainties about AI and geopolitics.
- Third, premature commitment to any one strategy forecloses others, undermining the capabilities, relationships, and advantages on which they depend (we detail the mechanisms in the “Risks of Foreclosure” section). Yet deferring commitment forecloses strategies by default because the security architectures that underwrite them go unbuilt. The key is to understand which foundations must be built ahead of need and when commitment can no longer be deferred.

¹ We find these terms to be reasonably self-defining, and none of our recommendations hinge on precise definitions or the specific form that advanced AI might take. Nevertheless, we offer these working definitions in the spirit of establishing common ground on the scope of claims and recommendations. We define *superintelligence* as AI that vastly exceeds human performance across most or all intellectual domains. *Artificial general intelligence* (AGI) is AI that matches or exceeds human performance across a wide variety of intellectual domains. *Frontier AI* is whatever AI systems represent the state of the art at a given time. *Advanced AI*, as used throughout this paper, spans the progression from current frontier AI toward superintelligence and is not tied to any single threshold or paradigm.

- Fourth, switching strategies once a crisis is underway is slow and costly, and an abrupt, forced switch can be destabilizing; the capacity to commit must be built deliberately and in advance.
- Fifth, in a competitive environment, building and demonstrating investments in that security architecture inevitably sends signals to rivals; the United States should therefore signal deliberately, shaping rivals' assumptions and options while recognizing that some signals advantage the United States and others do not.

We therefore recommend a **Freedom of Action strategy**: Build and preserve freedom of action to secure geopolitical advantage and ensure humanity's survival with agency through the transition to superintelligence.

This paper proceeds as follows. First, we set out the shared diagnosis and prescription. Then, we describe the seven archetypal strategies and the five pivots that discriminate among them. Next, we develop the risks of foreclosure and the cost of transition. Following that, we turn to signaling in a competitive environment. We then build out the Freedom of Action strategy and its goal hierarchy, as well as examine the resources and trade-offs required by the strategy. Finally, we address critiques of the Freedom of Action strategy and conclude. A forthcoming technical report will present the comparative analysis of the seven archetypal strategies on which many of these claims are built.

Diagnosis and Prescription

Beneath the disagreement about which strategy to pursue lies a substantial common ground about the situation. The diagnosis, from our paper *Finding Common Ground in AGI Strategy Debates* (Predd et al., 2026), can be stated in four observations:

- Frontier AI continues to advance, with no clear signal of whether, or when, progress toward superintelligence will slow, halt, or accelerate further.
- The stakes include decisive geopolitical advantage in the AI era and humanity's survival and agency.² Pursuing the first might risk trading off against the second.
- The U.S. government and frontier AI companies are indispensable actors, yet neither has a settled strategy or a detailed plan to protect humanity and secure geopolitical advantage.
- Each new model release or deployment is a potential triggering event for a rival's response, an incident, or a shift in the assumptions on which any strategy rests.

The prescription that follows from this diagnosis is that the United States must lead humanity to build two things: (1) the foundations for an ecosystem in which advanced AI and humans with agency

² Survival is more than biological persistence; it implies maintaining a viable population, genetic diversity, and technological capacity across millennia. Survival could be threatened by the misuse of AI by state or nonstate actors to create and employ new weapons of mass destruction or by an uncontrolled escalation to catastrophic war between great powers, among other scenarios. Preserving agency requires maintaining certain irreducible capacities: the ability to understand one's situation, the ability to deliberate and form intentions, the ability to act on those intentions, and the ability to revise course when circumstances change. Threats to agency include gradual disempowerment and concentration of power with authoritarian governments, among others. See Predd et al., 2026.

can coexist,³ and (2) a security architecture providing sufficient stability to build this ecosystem.⁴ But a sound strategy also requires a logic for effectively using the time available to decisionmakers in a setting in which institutionalized capacity for anticipatory policy is outstripped by the pace of technological progress. The seven strategies in the next section are competing answers about how to deliver the foundations for the ecosystem and security architecture described above, rather than disagreements about the goal.

Seven Archetypal Strategies

Seven archetypal strategies for the United States dominate current debates over how to navigate the path to advanced AI. The strategies divide into three families, and each strategy has a canonical statement in the public debate, a theory of victory, and a structural pathology.⁵

Coexistence Strategies

The coexistence family of strategies (A through C) accepts the prescription that the United States needs both an ecosystem in which humans and AI can coexist and a security architecture to preserve stability but differs on how to build it.⁶

A. Dominance. The United States leads the frontier and converts that lead into the leverage that sets the global order, accelerating development while suppressing rivals' programs through export controls, cyber operations, and supply-chain denial (Aschenbrenner, 2024). Even in success, a de facto monopoly concentrates power and risks locking in the dominant actor's values. Pluralism might persist among close partners but on terms that the dominant actor sets and can revoke.

B. Co-Development. Under the assumption that no actor can sustainably win a unilateral race, the United States leads a consortium (including China at a minimum) to co-develop safe superintelligence under shared governance, pooled compute, and layered verification and monitoring (Trager et al., 2023; Ho et al., 2023). Its pathology is twofold: (1) The founding members' values and internal power distribution become humanity's long-term default, potentially enshrining nondemocratic values, and (2) the consortium may prevent the worst outcome while becoming an arena for competition on everything it does not cover.

³ The ecosystem encompasses not just AI systems but also the digital, institutional, and physical environments within which they operate. The ecosystem's purpose is to create conditions for coexistence such that AI might help humans flourish according to their own values without replacing the sources of choice and meaning. See Predd et al., 2026.

⁴ We believe that any AI security architecture should create temporal space for the human-AI ecosystem to develop, prevent catastrophic conflict, establish minimum viable AI governance, and effectively manage AI incidents. See Predd et al., 2026. For more on national security challenges concerning AI, see Mitre and Predd, 2025.

⁵ A *pathology* is the way a strategy can degenerate even when functioning as designed.

⁶ We specifically chose *coexistence* because the often-overused term *alignment* does not describe the only paradigm that would enable humanity to survive with agency in a world of superintelligence, which might also be enabled by control mechanisms, architectural designs, symbiotic relationships, or yet undiscovered methods. Coexistence is intended to account for all of these paradigms and includes the concept of preserving human agency. Scenarios in which humans survive but lack agency would not qualify as coexistence in our model.

C. *Preparedness*. With neither decisive advantage nor formal cooperation achievable, the United States mobilizes a coexistence effort through informal coordination, including voluntary commitments, phased deployment, independent evaluation, incident reporting, and the domestic capacity to absorb and recover from AI incidents. This is reflected in such proposals as responsible-scaling policies (Dragan, King, and Dafoe, 2024; OpenAI, 2025, Anthropic, 2026) and defensive acceleration (Buterin, 2023). Its potential pathology is large-scale drift: *Preparedness* formalizes muddling through (Lindblom, 1959) as the de facto posture, normalizing accumulating damage and disempowerment.⁷

Denial Strategies

The denial family of strategies (D through F) accepts the prescription but, having assessed danger as close and coexistence infeasible, seeks to prevent or outlast advanced AI rather than build toward it.⁸

D. *Moratorium*. Judging catastrophic risk from the creation of advanced AI to be imminent—a position held by the Pause movement (Yudkowsky and Soares, 2025; Barnett and Scher, 2025; Future of Life Institute, 2023)—the United States pursues a verifiable global halt to AI development above specified thresholds, perhaps buying time for coexistence research to succeed.⁹ Its potential pathology is that the benefits of AI go unrealized, and the enforcement regime risks hardening into a global surveillance apparatus that curtails research and innovation and interferes with individual rights and liberties.

E. *Deterrence*. This approach holds that cooperative restraint is unachievable but suppression of superintelligence is feasible and is best articulated in the Mutually Assured AI Malfunction framework (Hendrycks, Schmidt, and Wang, 2025). On this approach, the United States halts its own frontier development above defined thresholds and uses the full toolkit of national power to prevent and, if necessary, destroy foreign programs before they reach dangerous capability, buying time for coexistence research to succeed. Its potential pathology is that the strategy buys survival at the price of a standoff that never settles into peace.

F. *Continuity of Society*. When every preventive strategy has failed or been judged infeasible, and coexistence is deemed impossible, the remaining objective is survival with some human populations retaining agency: The United States creates geographically distributed, biologically self-sustaining settlements hardened against AI-enabled threats (Ord, 2020). Even if successful at preserving humanity, many might suffer, and what remains of humanity could persist as a relic with little agency to meaningfully shape its own future.

⁷ The key distinction from the status quo or simple resilience strategies is the coexistence mobilization of *Preparedness*.

⁸ *Denial* is used here in this broad sense, halting, suppressing, or outlasting dangerous development, rather than in the narrower deterrence-by-denial sense of the strategy literature.

⁹ The feasibility of specifying thresholds, for any strategy, might itself be a challenge, particularly if capability gains are jagged and hard to observe.

Acceleration Strategy

G. *Acceleration*. This strategy judges that the dangers of constraining AI development exceed the dangers of AI development itself and that new security architectures will likely interfere with beneficial development. It is exemplified by effective accelerationism (Verdon, 2022) and the adjacent techno-optimism of Andreessen (2023). The United States relies on markets, competition, and rapid iteration to produce safety as a byproduct of useful AI, and removing regulatory friction promotes an abundance of material wealth. Its potential pathology is that abundance without strong governance leaves humanity exposed to harms that markets do not price, from pollution and the spread of misinformation to a deep dependence on systems it does not control.

Table 1 summarizes the seven strategies, with each strategy's canonical view, win state, and characteristic pathology.

Table 1. Seven Archetypal Strategies for Frontier AI

Strategy	Canonical View	Win State	Structural Pathology
Coexistence strategies (A through C)—judge coexistence feasible and build toward it			
<i>A. Dominance</i>	The Project (Aschenbrenner, 2024)	The United States reaches superintelligence first and suppresses rival efforts.	There is an extreme concentration of power and a lock-in of the dominant actor's values.
<i>B. Co-Development</i>	CERN for AI and International AI Organization (Hausenloy, Miotti, and Dennis, 2023; Trager et al., 2023)	The United States jointly governs superintelligence under shared control with rivals; no single actor is dominant.	Founding members' values and power become the long-term default, and the consortium could become a mechanism for intense competition.
<i>C. Preparedness</i>	Responsible scaling (Dragan, King, and Dafoe, 2024; OpenAI, 2025, Anthropic, 2026); decentralized, democratic, differential, and defensive acceleration (Buterin, 2023)	The United States mobilizes an informal coexistence effort while remaining resilient to accidents.	Muddling through normalizes accumulating damage and disempowerment.
Denial strategies (D through F)—judge coexistence infeasible (for now or forever); halt, prevent, or outlast advanced AI			
<i>D. Moratorium</i>	Pause (Yudkowsky and Soares, 2025; Barnett and Scher, 2025; Future of Life Institute, 2023)	The United States and rivals agree to halt capability development above agreed thresholds, which buys time for coexistence research.	Benefits are forgone, and enforcement hardens into a surveillance apparatus that infringes on rights and liberties.
<i>E. Deterrence</i>	Mutually Assured AI Malfunction (Hendrycks, Schmidt, and Wang, 2025)	The United States prevents the development of dangerous AI through threats and force if necessary.	Deterrence buys survival at the price of a standoff that never settles into peace.

Strategy	Canonical View	Win State	Structural Pathology
<i>F. Continuity of Society</i>	Existential security (Ord, 2020)	The United States creates self-sustaining human populations that can survive dangerous superintelligence.	Humanity persists as a relic, with limited leverage over its own future.
Acceleration (G)—judges the dangers of constraint to exceed the dangers of development			
<i>G. Acceleration</i>	Effective accelerationism (Verdon, 2022); techno-optimism (Andreessen, 2023)	The United States maximizes development and diffusion to capture maximum benefits and innovate around disruptions.	Abundance without governance leaves humanity exposed to harms that markets do not price, from pollution and misinformation to dependence on systems that humans do not control.

NOTE: CERN = European Organization for Nuclear Research (Conseil Européen pour la Recherche Nucléaire).

Five Pivotal Uncertainties

Which strategy a decisionmaker favors depends on judgments about five key uncertainties. Each is a question about the world whose resolution would favor some strategies over others. Reasonable experts disagree on these factors, and the disagreements cannot be conclusively resolved with currently available evidence.

Danger Proximity

How close is the danger from advanced AI? *Danger* here means wide-scale harm to humanity's survival and agency, including mass-casualty misuse, cascading societal harms, and AI-fueled or AI-incentivized escalation to catastrophic conflict and loss of control. Frontier capabilities are improving rapidly, but expert views vary on whether—and if so, when—these systems might pose such risks. If danger is close, strategies that buy time, build defenses, and invest in resiliency should be favored. If danger is not close or if humans are in greater danger without progress on advanced AI, the case for accelerated development strengthens.

Coexistence Feasibility

Can humans and advanced AI coexist in ways that enable human survival with agency? Coexistence here means a stable condition in which advanced AI operates at scale while humans retain the capacities that constitute agency. If coexistence is feasible, strategies that build and deploy advanced AI become more desirable. If coexistence is infeasible, the case shifts toward restraining or suppressing development or other means of protecting human civilization.

Restraint Feasibility

Can nations and firms cooperate to restrain development of AI? This question turns on whether verification of agreements is possible, the durability of cooperation under competitive pressure, and

the willingness of governments and companies to subordinate near-term advantage to long-term stability. If restraint is feasible, cooperative strategies become available. If it is not, restraint-based approaches are reachable only through more-coercive enforcement.

Decisive Strategic Advantage

Can any single actor achieve and maintain a decisive strategic or military and economic lead over its competitors (Bostrom, 2014)? Advantage runs along a continuum, from margins that merely favor the leader to a lead too decisive to contest; the pivot asks about the far end, whether a lead can be made decisive and durable. This depends on the scope of advantages that compound from a lead and the actor's ability to translate technical advantage into geopolitical leverage. If decisive advantage is achievable and the United States assesses that it can maintain dominance, a unilateral strategy may be sensible. If a decisive advantage is impossible, because of proliferation or other causes, no single actor can leverage AI to dictate global terms.

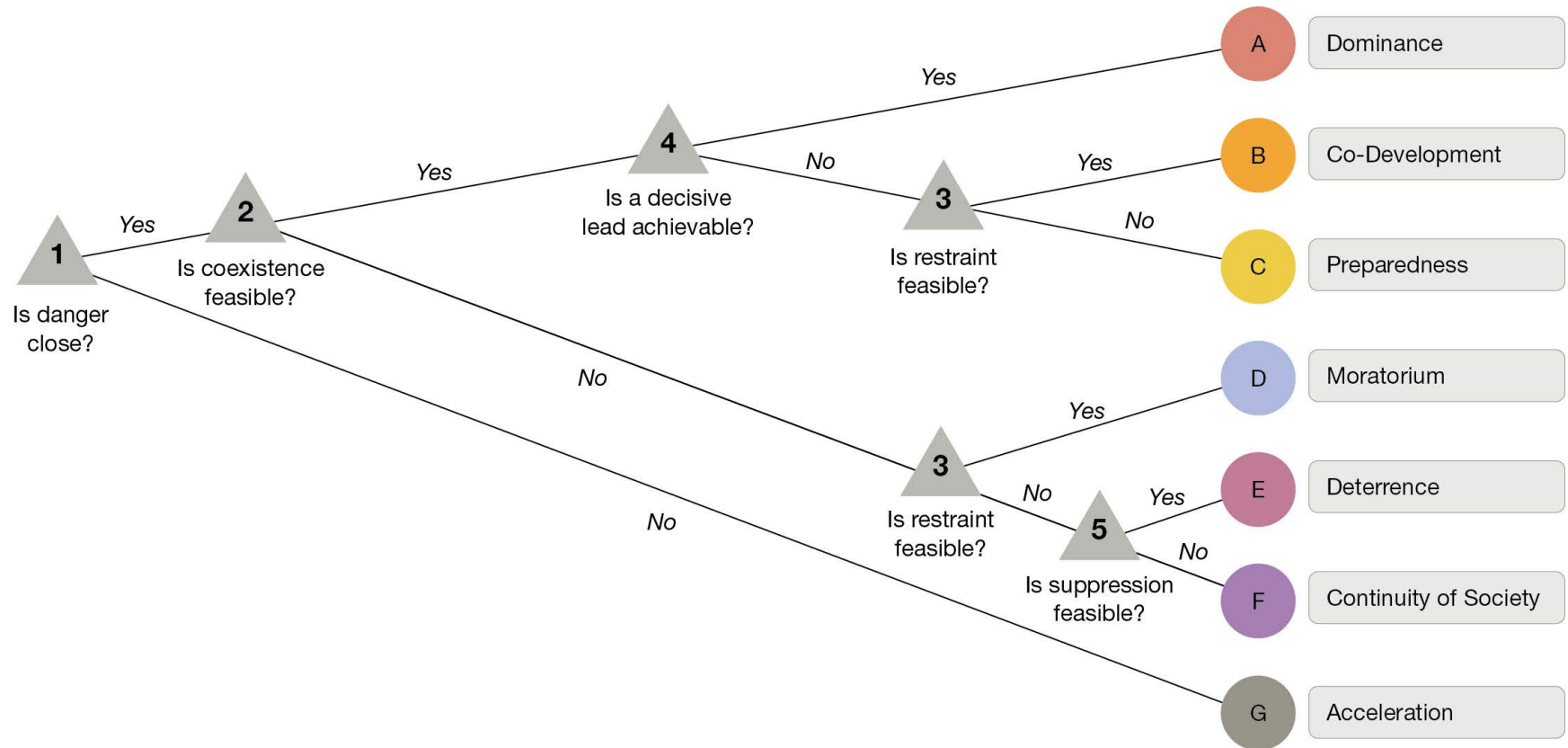
Suppression Feasibility

If cooperation fails, can the development of domestic or rival AI programs be suppressed by technology controls, deterrence, use of force, or other means? Suppression may be direct, degrading a program itself, or indirect, using economic and military pressure to compel a halt; the pivot turns on whether some combination can reliably stop development. This depends on relationships between the suppressor's ability to detect, attribute, and disrupt frontier AI development and the developers' ability to evade. If suppression is feasible, *Deterrence* becomes possible. If it is not, actors can press ahead with dangerous development even when others restrain, without fear of being restrained themselves.

These five questions do not resolve symmetrically. For four of them, evidence can confirm one answer but not the opposite. For instance, danger can be shown to be close but not shown to be far; coexistence can be demonstrated to work but not shown to be permanently impossible; a decisive advantage can be demonstrated but its impossibility never proven; and restraint can be observed to hold, although a record of broken agreements never proves future cooperation impossible. Suppression feasibility runs the other way, in that the creation of advanced AI can show that suppression does not work, but no success proves that it will work reliably against the next actor. For much of the strategic space, then, waiting for conclusive evidence means waiting for a proof that will never arrive.

Figure 1 maps how combinations of judgments on these factors lead to the seven strategies. Although drawn as a branching sequence for legibility, it should be read as a map of disagreements rather than a prescribed order of choices: For any two strategies, the diagram identifies the factors on which they most sharply diverge.

Figure 1. How the Five Pivots Resolve to Seven Strategies



Risks of Foreclosure

Each of the seven strategies depends on judgments about the pivots that current evidence cannot resolve. Committing to one strategy can foreclose others. *Foreclosure* occurs when an action or development removes a precondition required by another strategy for long enough that the strategy cannot be reconstituted within the window in which a decision would call for it. Some developments might even foreclose strategies in ways that are irreversible.¹⁰

We identify five practical mechanisms that foreclose strategies, differing in ways that bear on what to do about them, including what strategies they foreclose, how fast a foreclosed option could be rebuilt, and whether the foreclosure can be seen while it is happening.

Institutional Decay

Institutional decay forecloses strategies through the atrophy of expertise, organizational knowledge, or necessary infrastructure and state capacity. For instance, *Co-Development* and *Moratorium* depend on trust, diplomatic channels, and novel verification arrangements that have not been built or that weaken when they are not maintained. These strategic options can be lost with no decision to abandon them. Institutional decay can foreclose *Acceleration* as well, such as if the permitting, energy, and regulatory capacity that rapid deployment requires is not built. Some institutional foreclosure is abrupt and easy to date, such as withdrawal from an international engagement, while others may be due to gradual decay. An effort driven by a crisis might not be sufficient; for instance, the speed of any coordinated response to a pandemic depends on surveillance networks, laboratories, and legal authorities built well in advance. The U.S. withdrawal from the Open Skies Treaty in 2020 shows the abrupt form of this loss, in which an inspection regime and overflight expertise were dissolved by a decision, and neither could be easily restored.

Deployment Lock-In

Deployment lock-in forecloses strategies by embedding technologies so deeply that the political economy of reversal is prohibitive. As AI is embedded in social, economic, or institutional life, a halt can become costly. Lock-in and path dependence can foreclose *Moratorium* by creating powerful political and economic forces that oppose any slowdown. Lock-in can also narrow the domestic governance leverage that *Co-Development* and *Deterrence* would need to channel frontier development into international channels or suppress superintelligence outright. Accumulating dependence on AI weakens *Continuity of Society*, the strategy meant to remain available after the others fail. Each deployment decision is reasonable on its own and reversible when made, but the aggregate becomes prohibitive to unwind, and the loss is clear only once dependence is deep. Carbon-based energy infrastructure is one example of this mechanism; a century of individually reasonable deployment

¹⁰ The value of keeping options open when actions are irreversible and information improves is discussed as a “quasi-option value” in Arrow and Fisher (1974); see also Dixit and Pindyck, 1994.

decisions produced a political economy in which even broad agreement on harm cannot quickly unwind the dependence (Unruh, 2000).

Geopolitical Fait Accompli

Geopolitical fait accompli forecloses options from the outside. Suppose China parlayed an AI lead into a fait accompli, seizing Taiwan behind AI-enabled targeting and autonomous systems or successfully exporting its models so that they become the installed base across most of the world. In that case, the option to contest China might exist on paper, but leaders might consider the price too high in practice. Every strategy premised on contesting the Chinese position is foreclosed. *Dominance* is the clearest casualty because its win condition is exclusive AI leadership. Once a decisive lead is entrenched in Beijing, a U.S. bid for dominance is moot, and the remaining strategies must be entered late against an adversary that has already mobilized.

In principle, a rival's success can make any of the seven strategies moot. This pathway differs from diffusion and decay in that it is inflicted by another state's success. Keeping U.S. options open offers only partial protection against this pathway. The Soviet nuclear test of August 1949 illustrates the mechanism. The Baruch Plan for international control of atomic energy rested on the premise of a temporary U.S. monopoly. Although the United States was already facing diplomatic challenges, the Soviet detonation completely foreclosed the opportunity for a nuclear monopoly.

Agency Erosion

Agency erosion undermines the capacity to exercise any option at all. Every strategy presupposes a human capacity to choose and to act on the choice. That capacity can atrophy on its own, through dependence, cognitive offloading, and the accumulated weight of decisions ceded to AI systems, so that an option remains formally open while the practical ability to exercise it drains away (Moon and Boudreaux, 2026; Kulveit et al., 2025). Although there is some burgeoning evidence of these effects at the individual level (Sharma et. al. 2026), the collective, strategic form of this erosion (such as a polity losing the capacity to choose and execute a strategy) has not yet been observed. Measuring the impact of technology on human agency is an open research problem. This is why detecting this erosion before it nears the point of foreclosure is a core task of the monitoring function (discussed later as Objective 5).

Epistemic Foreclosure

Epistemic foreclosure can reduce options through the diffusion of knowledge. Once model weights, training methods, or specialized talent spread beyond a controlled boundary, the exclusive lead that *Dominance* requires might not be recoverable, and the leverage that *Dominance*, *Co-Development*, *Moratorium*, and *Deterrence* draw from a controllable frontier erodes. This foreclosure can happen quickly and visibly (because an open-weight release or a major leak is a dateable event), and it is effectively irreversible.

Encryption offers a partial analogue. For roughly two decades, the United States classified encryption products as munitions and subjected strong encryption to national security review and export restriction in order to limit the international spread of a capability that threatened signals intelligence collection. The controls did slow general dissemination, but they did not stop determined actors, and the regime was ultimately dismantled when there was no longer a public-private consensus. Commercial demand grew and domestic legal challenges narrowed the enforcement basis. By 2000, the restrictions had been substantially liberalized. But the underlying knowledge could not be recalled, so control had to be sought through other instruments. Vermeer (2024) draws a lesson for AI: Governance that depends on controlling the distribution of nonphysical assets, such as software or model weights, is the hardest kind to sustain.

Diffusion also carries gains, of course, such as the spread of U.S. technology and standards and widening the security-research base (Kapoor et al., 2024; White House, 2025), so its net effect on the strategy space is contested. Our foreclosure claim is narrower here, specifically that whatever diffusion's benefits, an exclusive lead (once diffused) might not come back.

Premature Commitment and Deferral Both Foreclose

The apparent choice between committing now and waiting passively for evidence to resolve the five pivots misstates the decision at hand. Both premature and deferred commitment can foreclose through the mechanisms above, a version of the dilemma that Collingridge (1980) identified for the control of any emerging technology.

Premature commitment triggers foreclosure through action. Commit to *Dominance* and the cooperative channels, verification arrangements, and shared evaluation frameworks that *Co-Development* and *Moratorium* require might atrophy, because the logic of racing to a decisive lead deprioritizes the institutions that those strategies would need. Similarly, a commitment to *Moratorium* might forgo the economic and scientific benefits of continued AI progress. Commit to *Acceleration* and deployment lock-in deepens as AI embeds in economic and institutional life faster than governance arrangements can form around it, foreclosing the strategies that depend on a governable frontier.

Deferral can trigger foreclosure through inaction. Several strategies need long-lead capabilities, including novel verification technology of compute, diplomatic channels, evaluation institutions, and a deep coexistence research base; these can erode through institutional decay while decisionmakers wait for conclusive evidence. If a government defers engaging with the pivot factors while deployment continues, it will effectively be executing *Acceleration* in practice. Meanwhile, deployment lock-in advances, and if monitoring lapses, the government loses the ability to assess the pivots that any future strategic choice would require.

To avoid the risks of strategic foreclosure, the United States should actively build and maintain the foundations across the strategic space, commit when the pivot factor evidence warrants, and intentionally make decisions to preserve or abandon capabilities that might foreclose options if not built in time. We later develop this course as the Freedom of Action strategy.

AI Strategy in a Competitive Environment

Much of the preceding analysis of the risks of foreclosure is framed in a single-actor setting, as though the United States can make strategic choices independently. Yet, the United States' rivals and partners face their own strategic choices about how to manage frontier AI on the uncertain path to superintelligence, and those choices will shape what is feasible and sensible for the United States. Indeed, in a competitive environment, an AI strategy is a signal as much as a framework for action, revealing what the state believes about how fast capabilities will advance, whether advantage will prove decisive, or whether the adversary can be lived with or must be overcome. In response, rivals will read and adjust, invest, hedge, accelerate, or restrain against the future that they infer the United States is planning for and send signals of their own that Washington must read in turn.

Commitment to a particular strategy may be the loudest such signal but far from the only one. Every investment and demonstration in a security architecture tells rivals something about what the United States can do and might intend, perhaps even before a strategy is explicitly chosen. Silence and ambiguity are read too, so the real question is what to signal, to whom, and when. Some signals work in the United States' favor; for instance, a demonstrated ability to detect and defeat a covert program can deter one from ever being launched. But some carry unintended signals alongside the intended one. A visible move toward suppression capability, meant to deter, can also signal that the window for development is closing, giving a rival a stronger reason to accelerate before the window shuts. Other signals demonstrate openness, such as sharing the results of a safety evaluation or accepting a reciprocal inspection, showing that cooperation is still on the table. Rivals are signaling back the whole time, so the United States is always reading as much as it is read.

Two implications follow. First, the United States should choose signals for their effect on the strategy space and favor those that widen U.S. advantage and freedom of action or move a rival toward restraint, as well as weigh carefully any that would foreclose U.S. options or harden a rival's resolve. For illustration, demonstrating a credible independent evaluation and verification regime widens the availability of possible strategies but would not seem to foreclose any. This regime serves both deterrence, by enhancing awareness of capabilities, and cooperation, by offering a template for reciprocal inspection. Alternatively, letting evaluation and verification resources and talent visibly dissipate narrows the availability of several strategies.

Ambiguity also has costs. Allies who cannot read U.S. intent cannot align with it, and a rival who reads hesitation rather than deliberation might be emboldened to attempt the *fait accompli* that we identify in the previous section as a foreclosure mechanism.

The second implication is that the United States should manage what it withholds as deliberately as what it reveals. Deterrence requires demonstrating capability, but revealing intent is a different calculation, because specifying the conditions under which a capability would be used might allow a rival to plan accordingly. For instance, U.S. policy toward Taiwan has long exploited this asymmetry; carrier deployments and arms sales establish that the United States can intervene in a cross-strait conflict, while official silence on what the United States would actually do means that Chinese planners have to consider multiple contingencies.

Offensive counter-AI capabilities are another example. Suppose that the United States develops cyber tools capable of corrupting a rival's training data or halting a frontier training run. This might serve multiple strategies, including *Dominance* and *Deterrence*. Disclosing that the tools exist may

strengthen these strategies, but disclosing their explicit purpose does not. Stating which purpose takes precedence would tell Beijing what to harden or perhaps how to preempt. Ambiguity about intent preserves the U.S. option to decide later.

A Strategy to Maintain Freedom of Action

The previous sections make the point that the United States needs both an ecosystem in which humans and AI can coexist and a security architecture that keeps the path to it stable, but the seven strategies that would deliver that prescription cannot be objectively ranked on the evidence available, and that evidence may not resolve on the schedule that decisions demand. Committing early or deferring decisions forecloses strategic alternatives, while switching once a crisis is underway is too slow and costly and, in some cases, infeasible. Moreover, in a competitive contest, any commitment inevitably affects competitive dynamics and a rival's strategic choices.

These conclusions suggest that the United States should build and preserve options across the strategic space, make the investments that keep perishable options alive, signal to hold deterrence and cooperation credibly open at once, and keep ready the capacity to seize advantage when windows open and to commit when the evidence clearly favors one strategy. We call this the *Freedom of Action strategy*.¹¹

Overarching Objective

The objective of the Freedom of Action strategy is to build and preserve the ability to secure U.S. geopolitical advantage and ensure humanity's survival with agency through the transition to superintelligence. This involves the capacity to seize advantage as windows of opportunity open and to pursue whichever of the seven strategies the uncertainties come to favor.

The objective requires us to differentiate optionality as ends versus means. Strategic optionality, keeping all seven strategies available to the United States, is instrumental at this present moment (mid-2026) to preserve humanity's survival with agency. However, humanity's agency—the capacity of people to understand their situation, decide, and act on the decision—is an end in itself or an inherent good that is important from a broad set of moral and religious views. Preserving agency also serves U.S. geopolitical interests (Cibralic, 2026). Our claim in this paper is that, under present uncertainty, preserving strategic optionality is the most reliable means of securing geopolitical advantage and human survival with agency.¹² Freedom of action is therefore the guiding policy, an overall approach for coping with the diagnosed obstacle of abundant uncertainty (Rumelt, 2011).

¹¹ This approach is closest to what Winter and Bullock (2026) call “radical optionality,” a domestic-governance strategy of preserving democratic decisionmakers’ ability to respond competently as circumstances evolve; the strategy developed here extends that logic to the geopolitical contest, to the structured space of archetypal strategies and pivots, and to the mechanisms by which options are foreclosed.

¹² Several of the seven strategies can be read as bets on how this tension should be resolved. *Dominance* bets that geopolitical advantage is the surest route to survival with agency, reasoning that whoever reaches superintelligence first can set the terms for successful coexistence. *Moratorium* bets the other way, subordinating advantage to survival with agency and accepting forgone

The Freedom of Action Strategy's win state is that the United States preserves the capacity to execute whichever of the seven strategies the evidence comes to favor and executes it in time. Its structural pathology, discussed below, is that the apparatus built to preserve options acquires a constituency for keeping them open and resists commitment after the evidence has resolved. The strategy requires the United States to both invest in building new capabilities and capacities and adapt existing ones at the same time that it operationalizes them to navigate a fast-paced and dynamic technological and geopolitical environment.

Investments: What the United States Must Build

The United States must build a human-AI ecosystem (Objective 1), the elements of an AI security architecture (Objective 2), national security enterprises fit for purpose for the AI era (Objective 3), and the capacity of citizens, firms, and governments to respond (Objective 4). These are near-term investments worth making across a wide variety of futures, and many carry lead times that require they be built now.

Objective 1: Build the Human-AI Ecosystem.

This objective is the affirmative half of what the strategy exists to secure, a future in which humans and advanced AI coexist with human agency intact. It is also the foundation of the security architecture among geopolitical actors because durable geopolitical advantage in the AI era depends on getting this ecosystem right. Its lead times may be long.

- **Advance and sustain the AI safety and coexistence research base.** Fund the science that makes advanced AI steerable, bounded, and able to be shut down, and keep many independent lines of inquiry alive, at levels that track capability (on control, see Greenblatt et al., 2024; Korbak et al., 2025). This base is the scientific foundation for whether coexistence proves feasible. Letting this decay or failing to build closes the entire coexistence branch of the strategy space and eliminates conditions for pivoting away from *Moratorium* and *Deterrence*.
- **Build tools and institutions that preserve human agency.** Develop the interfaces, oversight bodies, and authorities that keep people deciding the consequential questions—such as those with large or hard-to-reverse stakes—as more-routine choices are delegated to AI systems. Deciding which questions count as consequential is itself a governance choice. These decisions should be assigned to accountable institutions, made in advance, and revisited as delegation spreads rather than left to default trajectories.
- **Establish arrangements for broad benefit and legitimacy.** Set the terms that earn and sustain the public consent on which a sustainable strategy depends. This includes access to advanced capabilities, accountability for harms, and a broad distribution of AI-driven economic gains through such instruments as competition policy, public-benefit requirements, and workforce transition support.

benefits and a frozen relative position as the price of buying time. *Continuity of Society* abandons geopolitical advantage to focus on survival with agency.

- **Prepare society for AI-driven disruption.** Build the economic transition capacity, information-integrity infrastructure, democratic resilience, and governance of AI-driven dependency that keep the ecosystem's social foundations intact as AI capabilities advance.
- **Shape the incentives and structure of the AI ecosystem.** Create indicators and tools to align market incentives with safety, dampen racing dynamics that reward dangerous corner-cutting, set standards that become global defaults, manage AI agent ecosystems, and monitor power concentration.

Objective 2: Build the AI Security Architecture

The AI security architecture is the set of technical and institutional capabilities that let the United States advance and secure its own AI development progress, detect dangerous AI developments, verify others' behavior, contain incidents, manage geopolitical and technological risk, and deny capability to reckless actors. Illustrative examples of these domains are provided below. Many of these capabilities also contribute to the development of the human-AI ecosystem in Objective 1. This highlights two unusual features of any security architecture for frontier AI: (1) The ecosystem itself performs security functions, because a healthy human-AI ecosystem reduces the very risks the architecture exists to manage, and (2) much of the security architecture depends on the private sector and is carried out through partnerships with AI companies and the broader technology and research community. This includes multiple capabilities:

- **Situational awareness.** Stand up independent evaluations of domestic and international frontier-model capabilities and dangerous-capability thresholds, detection of large-scale training runs, and monitoring of compute infrastructure and algorithmic progress.
- **Compute visibility and governance.** Develop the means to know where frontier-scale compute sits and how it is used, and the levers to govern it. Compute governance serves every strategy differently, and its utility depends on the assumption that frontier development requires large, identifiable compute concentrations.
- **Incident detection, containment, and response.** Build the capacity to catch, isolate, and recover from AI incidents before they cascade, including targeted infrastructure shutdown and continuity-of-government planning.
- **Verification, monitoring, and cooperative infrastructure.** Fund the technology to verify potential agreements and maintain the standing diplomatic channels and technical working groups through which rivals and allies are engaged.
- **Deterrence and enforcement.** If strategies require it, build the ability to suppress dangerous AI initiatives and enforce moratoriums.
- **Counter-AI and resilience.** Develop adversarial techniques against AI systems, counter-autonomous-system capabilities, AI-enabled cyber defense, and frontier-model weight and data center security.
- **Domestic governance and enforcement capability.** Build the regulatory expertise, independent evaluation infrastructure, and standing channels to frontier companies through which the other capabilities in this objective are authorized and applied, and keep registration, licensing, and emergency authorities ready rather than improvised.

Objective 3: Adapt the Legacy National Security Architectures and Institutions

The institutions and capabilities that already guard the United States—including the nuclear enterprise, biosecurity, the conventional military, the intelligence community, counterterrorism, and cybersecurity—must each adapt for the AI era that is emerging. Diplomacy and economic statecraft require the same adaptation, and the national security enterprise as a whole might need to be reconceived in both mission and operations. This work must run in parallel and begin before AI outruns the institutions built for an earlier era.

The adaptations run enterprise by enterprise. For instance, the nuclear enterprise adapts deterrence, arms control, counterproliferation, and surety (across doctrine, concepts, and capabilities) for the AI era; biosecurity scales for AI-enabled biological risk; the conventional military adapts forces, doctrine, and resilience for AI-enabled conflict; the intelligence community adapts collection, analysis, attribution, and covert action for the AI frontier; counterterrorism adjusts to the changed scope and scale of AI-enabled nonstate actors; and cybersecurity adopts AI tools, defends critical systems against AI-enabled offense, and hardens U.S. AI systems against adversarial attacks.

Objective 4: Build the Capacity of Citizens, Firms, and Governments to Respond

This objective builds the human and institutional readiness that the strategy will draw on. This includes the leadership, beneficial uses, crisis machinery, and literacy without which even a correct strategy cannot be executed in time. The foundations of this competition are as much social and cultural as technological. No single actor can build everything needed for readiness.

- **Prepare leaders.** Equip decisionmakers with the awareness, networks, expertise, and rehearsed experience to act well in fast-moving AI crises. Leaders who first encounter these dilemmas in the situation room are already behind. Preparation cannot start at the top: The staffs, commands, and operating levels that frame options for principals need the same rehearsed familiarity earlier, because leaders act on the choices their institutions surface.
- **Cultivate safe, positive uses.** Develop and diffuse the beneficial applications, in security, prosperity, and human flourishing, that any committed strategy will ultimately need to deliver and that sustain the political support necessary for a sustainable strategy.
- **Build break-glass plans and capabilities, and institutionalize AI crisis planning.** Create contingency plans and the means to execute them and make crisis planning a standing function rather than an improvisation, naming the contingencies in advance.
- **Raise AI literacy.** Broaden public and institutional understanding of AI's capabilities, limits, and risks so that decisions under pressure rest on informed judgment rather than hype or fear.

Sequencing the Investments

Freedom of action might be misread as emphasizing waiting. But it is the opposite: An option exists only when it has been built, and most of the capabilities that Objectives 1 through 4 require take substantial investment in time and resources. The strategy therefore involves significant front-loaded effort now to preserve optionality later.

The order of the needed work depends on the time required to build a capability (lead time) and the window in which it will be needed (need horizon). Where the lead time is longer than the need horizon, the capability must be started immediately or written off as unavailable. The Freedom of Action strategy would make these choices explicit instead of allowing options to expire silently.¹³

Lead times differ by an order of magnitude. Institutions and trained people take years to a decade, verification and monitoring infrastructure is comparatively slow, energy and compute run on multiyear cycles, and some evaluation and software capabilities take only months. In this way, the Freedom of Action strategy provides a construction schedule for necessary capabilities.

Operations: What the United States Must Do

The United States must put these capabilities to use by watching for indications and warnings, guarding against catastrophe and lock-in, and seizing advantages as opportunities present.

Objective 5: Monitor for Indications and Warnings

Monitoring is the standing capacity to see clearly enough to know which strategies remain open, which are closing, and when the evidence has moved enough to act.¹⁴ One object of monitoring is the five pivots. The United States would seek out indications of danger proximity, coexistence feasibility, restraint feasibility, decisive advantage, and suppression feasibility, together with the rundown of each strategy's decision clock. The warnings issued are that a strategy is closing, that the evidence has begun to favor one strategy, or that a latest decision date is near. A concrete indicator list is follow-on work to this paper, and portions of it will likely need to be developed in classified channels. Different technological pathways to superintelligence, from continued scaling to algorithmic breakthroughs to the emergent coordination of many AI systems, would produce different signatures, and some might produce little advance warning at all. The monitoring function must therefore watch not only for incremental advances along the current trajectory but for paradigm shifts and emergent dynamics. Monitoring is run as a strategic-warning function rather than a periodic assessment on the model of the enterprises built to detect strategic surprise: watching the frontier, watching the competition, and putting warning in front of decisionmakers in time to act.

- **Anticipate frontier advances.** Maintain a rolling horizon scan of capabilities likely to cross significant thresholds within the next several months to years, through sustained engagement with frontier companies and other relevant sectors.
- **Identify militarily significant developments.** Run structured capability sprints and military-relevant benchmarks to determine when an advance is decision-relevant.

¹³ Note that crash programs can compress timelines but only by so much and might have vastly different outcomes when applied to engineering problems rather than to building institutions or diplomatic dialogues.

¹⁴ The monitor-and-commit structure of this strategy draws on the literature on decisionmaking under deep uncertainty and adaptive planning (Lempert, Popper, and Bankes, 2003; Haasnoot et al., 2013; Marchau et al., 2019). The operational picture is also informed by RAND's Day After Artificial General Intelligence exercises (Predd, Chesson, and Smith, 2025; Smith et al., 2025; Smith et al., 2026).

- **Watch and red-team rivals.** Track rivals' programs, deployments, and intent—especially China's AI stack and People's Liberation Army adoption—and rerun the red assessment with each model generation.
- **Set indicators and warning thresholds.** Specify in advance the observable signals that would move an informed assessment, such as signals related to frontier development, rivals' deployments, and the pivotal assumptions behind each strategy. From these signals, maintain an updated read of which strategies the evidence leaves open, which it has begun to favor, and how much time each decision clock leaves.
- **Drive decisions.** Get the warning and the decision it calls for to the right decisionmaker while there is still time to act. The history of strategic warning cautions that the function fails at the receiving end as often as the producing end, when the recipient misreads the signal or is unprepared to act on it (Wohlstetter, 1962; Hoehn and Shanker, 2023); warning can only warn, so the measure of success is whether the United States decides and acts.

Objective 6: Guard Against Catastrophe and Lock-In

This objective protects the strategy space while the evidence comes in, buying time and ensuring that a catastrophe or a *fait accompli* does not settle the outcome before the evidence arrives. Buying more time keeps every other option open longer.

- **Deny, disrupt, and degrade reckless advancement.** Slow the fastest and least cautious actors. *Reckless* here means development or release that crosses the dangerous-capability thresholds of Objective 2 without corresponding safeguards. This might leverage instruments discussed previously, such as export controls, licensing mechanisms, and enforcement capabilities.
- **Manage dangerous proliferation.** Apply dual-use principles to manage the spread of AI capabilities whose release would foreclose strategies that depend on a governable frontier or provide meaningful uplift toward catastrophic harm. In practice, this means prerelease review against defined thresholds, staged release with monitoring, and export control–style licensing. All of these are admittedly imperfect instruments.
- **Deter opportunistic aggression by AI-enabled adversaries.** Deter adversary aggression, including aggression by actors intent on preventing or protecting AI progress outside an agreed-upon governance framework or empowered by a temporary first-mover advantage. The instruments are those of Objective 3 adapted for the AI era: conventional and nuclear posture, cyber response options, and declaratory policy.
- **Hold strategic stability.** Maintain the conditions, including crisis stability and alliance cohesion, under which any AI governance regime can function. The means leveraging the architecture of Objective 2: crisis-communication channels, transparency and confidence-building measures, and alliance consultation on AI issues.
- **Keep a nuclear firebreak.** Do not integrate AI into nuclear command, control, and communications without demonstrated safeguards. This development could foreclose every strategy at once, so the firebreak must hold whichever strategy is in force.
- **Manage incidents.** Detect, attribute, contain, and recover from AI incidents before they cascade or eliminate strategic optionality.

- **Build and maintain resilience.** Sustain the resilience of critical infrastructure, institutions, and society against AI-enabled disruption so that the United States can absorb shocks and retain the capacity to choose.
- **Act on expiring options.** Track each strategy's latest decision date and treat its arrival as a decision to preserve the option or accept its foreclosure. Options are also lost without any decision at all. Capabilities, relationships, and expertise held in reserve must be exercised rather than shelved, and each deployment that deepens dependence should be weighed for what it will cost to unwind because the individual choices are reasonable and reversible, while the aggregate is neither.

Objective 7: Seize Advantage and Signal Deliberately

Windows of advantage might open unpredictably, and how advantages are seized is itself a signal that will be read by rivals. Each move must be chosen for the advantage it secures and the message it sends.

- **Extend the window.** Take the actions that keep a favorable opportunity open longer by sustaining a lead or slowing diffusion, so that there is room to use the opportunity deliberately.
- **Exploit and integrate.** Turn the advance into operational advantage quickly and push it into the decision cycles, concepts, and forces that can use it before it diffuses. Seek gains that leave the United States with more choices afterward unless a commitment that narrows options would provide a long-term sustainable advantage.
- **Signal.** Treat each demonstration, disclosure, and silence as a deliberate message: Demonstrate selectively to show that aggression could be deterred or defeated, share selectively to keep cooperation credible, and hold the rest in reserve. Because much of the capability being signaled sits with private companies rather than the state, effective signaling requires close government-industry coordination over what is demonstrated and disclosed, run through standing mechanisms rather than ad hoc requests: cleared-industry forums on the defense-industrial model, pre-negotiated disclosure agreements, and crisis playbooks exercised jointly under Objective 4. Keep deterrence and cooperation visibly open at once, so a rival can assume neither a wall nor a hand. Capability-sharing is itself a strategic instrument (a question of what to share, with whom, and when), and signaling is also how the United States dampens a rival's incentive to race.
- **Commit, by evidence or by the clock.** The evidence and the clock govern two different decisions. When the evidence on a pivot becomes compelling, commit to the strategy it favors on the merits.¹⁵ Short of that, each perishable capability forces a narrower decision as its latest start date approaches: Begin the long-lead work now and keep the strategy reachable, or let the option close. Treat each latest decision date as a deliberate choice to preserve an option or accept its foreclosure instead of letting the choice be made by default.

¹⁵ Note that this logic works for all strategies except *Deterrence*. As described in the "Five Pivotal Uncertainties" section, the availability of suppression can never be conclusively proven with evidence. *Deterrence* must be chosen based on a judgment that the danger is bad enough to act and that the downside risks, such as escalation, are manageable.

Figure 2 summarizes the Freedom of Action strategy and the seven objectives.

Figure 2. A Freedom of Action Strategy



Resources and Trade-Offs

A successful strategy needs to be tested against the availability of resources, and admittedly we do not assess costs in detail in this paper. But we do seek to provide an order-of-magnitude reasoning behind that accounting, including how large the commitment plausibly is, where the binding constraints lie, and which trade-offs the strategy forces.

To be clear, the commitment for the United States is a complete societal and institutional reorientation rather than a single program. The United States has made reorientations of this class at least twice in the past century. The early Cold War era produced the National Security Act of 1947 and with it the Department of Defense, the Central Intelligence Agency, and the National Security Council. NSC-68 and the Korean War then roughly tripled defense spending within two years. A decade of sustained investment built the analytic, warning, and verification apparatus on which every later nuclear strategy depended (Gaddis, 2005). And later, after September 11, 2001, the Department of Homeland Security and the Office of the Director of National Intelligence were stood up within four years, homeland security spending ran on the order of a trillion dollars over the following two decades, and intelligence and law enforcement attention was reallocated (Hoehn and Shanker, 2023).

The record of those reorganizations is mixed, and the post-2001 case in particular may have led to mission sprawl and waste. But both cases show that the commitment is measured in societal and institutional attention sustained across administrations more than in any single budget line.

Every archetypal strategy requires a comparable reorientation of government and societal attention. Each strategy needs the monitoring apparatus, the security architecture, the adapted institutions, and the machinery to manage its own pathological mode of failure. The pathologies catalogued in the “Seven Archetypal Strategies” section are themselves resource sinks: a *Dominance* posture must oversee its own concentration of power just as a *Moratorium* posture must ensure democratic oversight of the surveillance capabilities needed to sustain enforcement. The seven strategies have many capabilities in common but also point resources in different directions. Nevertheless, this requires the United States to make orders of magnitude of commitment beyond the status quo. The cost differential among the archetypal strategies is dwarfed by the cost of implementing any serious strategy or the cost of strategic drift.

Given the possible extensive economic returns to AI, the binding constraint is plausibly something other than money. Private AI capital investment for capability development reportedly runs to hundreds of billions of dollars annually (Sajadieh et al., 2026). By reducing pressures to slow AI development until evidence warrants, the Freedom of Action strategy can help promote economically beneficial applications of AI, and in principle, the growth that the strategy protects can help defray its costs. The more practical scarcities that bind are political will, attention, institutional capacity, technical talent inside government, political bandwidth, and the public-private coordination on which most of the architecture depends.

There are nonetheless real trade-offs for the U.S. government, business, and civil society. Senior leader attention is finite, and this strategy would draw from other priorities. Much of the architecture also requires appropriations, and some of it might require authorities that do not now exist. And the strategy asks political leaders to sustain some investments whose payoff is the absence of foreclosure and maintaining optionality, which is a benefit that might be hard to calculate.

Critiques

The Freedom of Action strategy faces several critiques.

The Costs of—and Capacities to—Maintaining Optionality

Preserving optionality across a variety of futures for such a transformational technology is inherently expensive. The previous section argues that the scale of the demand might be comparable across strategies but more importantly is dwarfed by the gap between any strategy and the status quo. But there remains a critique that scarce budgets, attention, and institutional capacity determine which options are worth keeping open. Some strategic options should be allowed to lapse when their marginal value does not meet the marginal cost of keeping or building them.

To respond, we highlight three points. First, as noted in the previous section, the resource demands may be broadly comparable across the seven strategies and are dwarfed by the gap between adopting any strategy and remaining at the status quo, so the increment attributable to preserving optionality itself is fairly modest. Second, the costs do not all fall on the federal budget. Any serious response to the five pivots requires a whole-of-nation effort, and much of the necessary work—alignment, evaluation, security—is already being funded by firms and other actors. Third, the binding constraint may be not money but rather time and prioritization. The capabilities that unlock multiple strategies at once offer a natural focal point for scarce attention, and they give the dispersed actors who must contribute a common object to build toward. More generally, attention to transformational AI will likely not stay scarce. A technology of this consequence will captivate governments. The framework seeks to give that attention a constructive agenda organized around a strategy rather than improvised in a crisis.

The Shortest Timelines

The strategy presumes sufficient time to monitor developments and build options across strategies. But progress is accelerating, with frontier developers forecasting recursive self-improvement within a few years (Kokotajlo et al., 2025). If they are right, we may not have time for a coordinated strategy at all.

However, short timelines do not favor an early bet on a single strategy. The shorter the timeline, the less time there is to recover from a wrong choice, which raises the value of focusing resources on the shared foundations that pay off no matter which strategy proves correct. Short timelines sharpen the sequencing logic of the strategy development by shortening the latest decision dates for building capabilities. Short timelines also suggest a need for front-loading the capabilities with the longest lead times and widest cross-strategy value first at crash-program scale.

Of course, there might be exceptionally fast takeoffs, where no foundation can be built in time and no evidence will resolve before the choice is forced. We acknowledge that, in these cases, the Freedom of Action strategy might not have the room it needs to operate—which may be true for the other strategies as well. There might be no solution to the shortest timelines, but the best mitigation is to build the shared foundations now rather than defer and accumulate risk.

The Will to Commit

The Freedom of Action strategy focuses on the dangers of premature commitment to a specific strategy archetype. Actors who commit to a strategy without conclusive evidence of the resolution of pivot factors risk foreclosing strategic options when the pivot factors resolve. But another mode of failure may be the opposite: risk-averse institutions, slow political processes, and visible time pressure combining into too little, too late.

We proposed the latest-decision-date analysis to counter that possible drift, but a decision rule cannot by itself supply the will to act on it. The strategy's success presumes institutions prepared to commit when the evidence arrives and to decide deliberately rather than by default as the decision date to build needed capabilities passes. That capacity is built in Objective 4. Leaders and staff who have rehearsed a commitment decision and have a break-glass plan already drafted for the strategy that the evidence favors face a smaller burden than those improvising one under pressure. Rehearsal does not supply will. It removes the substitutes for will that can undermine good decisionmaking, such as options not being available, decision consequences examined under undue time pressure, or authority not residing in those who need to take action.

A Recipe for Chaos

Holding several strategies open invites the objection that a nation with a foot in seven canoes is headed for chaos rather than flexibility. But preparing for multiple contingencies without knowing which will arrive is smart practice of effective institutions. Consider the U.S. military's emphasis on force planning against multiple contingencies, in which combatant commands maintain standing plans, and equipment is prepositioned in theaters. The alternative for this type of planning under uncertainty—building the force for a single predicted war—is itself irresponsible. We do not suggest executing seven strategies at once but rather preparing and planning across the strategy space.

Freedom of Action's Own Structural Pathology

Finally, the strategy might itself foreclose options through its own apparatus. A standing monitoring, licensing, and review architecture is subject to the same deployment lock-in described in the "Risks of Foreclosure" section. Agencies and constituencies form around the strategy, and what was built to preserve options can harden into a permanent brake on the acceleration end of the strategy space. This suggests that the Freedom of Action architecture should be subject to the same discipline it applies elsewhere: periodic review keyed to the five pivots so that elements no longer providing optionality value are wound down rather than defended.

Conclusion

The possible emergence of superintelligence is a strategic challenge without precedent, and the temptation is to resolve its uncertainty by force of will, to pick the strategy that present instinct favors and commit to it. That is a risky choice that misunderstands the structure of the problem and the

extensive uncertainty we face. The opposite temptation to be sclerotic or paralyzed under uncertainty is not better.

The seven strategies cannot be conclusively ranked on the evidence available; commitment in either direction, acting now or waiting, forecloses options the United States will need; and a premature commitment also might signal to a watching rival how to beat it. The disciplined response is to manage the dilemma by building and preserving freedom of action, hold both deterrence and cooperation credibly open, seize advantage when windows of opportunity open, and commit when the case is compelling. Capabilities whose absence would foreclose options should be preserved or abandoned with intention. The capacity to choose, and the readiness to choose well when the moment comes, is how the United States secures both geopolitical advantage and humanity's survival with agency.

This is a demand for work that begins now, before any crisis arrives. It means building and deploying the capabilities that serve every strategy and beginning now to build those with lead times that exceed the warning the United States is likely to get. It means standing up the monitoring to read the five uncertainties and the decision machinery to act on them, demonstrating enough capability to shape rivals' choices without disclosing intent, and rehearsing the response. None of this can be improvised mid-crisis, and each capability left unbuilt already forecloses options. The burden falls on the U.S. government, the frontier companies, businesses, and civil society together, in service of the people whose agency the strategy ultimately exists to protect.

Bibliography

- Amodei, Dario, "Machines of Loving Grace: How AI Could Transform the World for the Better," DarioAmodei.com, October 2024.
- Andreessen, Marc, "The Techno-Optimist Manifesto," Andreessen Horowitz, October 16, 2023.
- Anthropic, "Responsible Scaling Policy," version 3.0, February 24, 2026.
- Arrow, Kenneth J., and Anthony C. Fisher, "Environmental Preservation, Uncertainty, and Irreversibility," *Quarterly Journal of Economics*, Vol. 88, No. 2, 1974.
- Aschenbrenner, Leopold, "Situational Awareness: The Decade Ahead," June 2024. As of August 31, 2026: <https://situational-awareness.ai/>
- Barnes, Beth, Hjalmar Wijk, and Lawrence Chan, "Responsible Scaling Policies (RSPs)," blog post, Model Evaluation and Threat Research, September 26, 2023.
- Barnett, Peter, and Aaron Scher, *AI Governance to Avoid Extinction: The Strategic Landscape and Actionable Research Questions*, Machine Intelligence Research Institute, May 2025.
- Baum, Seth D., David C. Denkenberger, and Jacob Haqq-Misra, "Isolated Refuges for Surviving Global Catastrophes," *Futures*, Vol. 72, September 2015.
- Bengio, Yoshua, Stephen Clare, Carina Prunkl, Malcolm Murray, Maksym Andriushchenko, Ben Bucknall, Rishi Bommasani, Stephen Casper, Tom Davidson, Raymond Douglas, et al., *International AI Safety Report 2026*, UK Department for Science, Innovation and Technology, DSIT 2026/001, February 2026.
- Bernardi, Jamie, Gabriel Mukobi, Hilary Greaves, Lennart Heim, and Markus Anderljung, "Societal Adaptation to Advanced AI," arXiv, arXiv:2405.10295v3, January 23, 2025.
- Bostrom, Nick, *Superintelligence: Paths, Dangers, Strategies*, Oxford University Press, 2014.
- Bostrom, Nick, "The Vulnerable World Hypothesis," *Global Policy*, Vol. 10, No. 4, November 2019.
- Buterin, Vitalik, "My Techno-Optimism," blog post, November 27, 2023.
- Chessen, Matt, and Swaptik Chowdhury, *Pivots and Pathways on the Road to Artificial General Intelligence Futures*, RAND Corporation, PE-A4178-1, October 2025. As of August 12, 2026: <https://www.rand.org/pubs/perspectives/PEA4178-1.html>
- Christiano, Paul, "What Failure Looks Like," AI Alignment Forum, 2019.
- Cibralic, Beba, *The Strategist's Dilemma: Agent-Neutral and Agent-Relative Reasons in U.S. AI Strategy*, RAND Corporation, PE-A4931-1, July 2026. As of August 20, 2026: <https://www.rand.org/pubs/perspectives/PEA4931-1.html>
- Collingridge, David, *The Social Control of Technology*, St. Martin's Press, 1980.
- Dixit, Avinash K., and Robert S. Pindyck, *Investment Under Uncertainty*, Princeton University Press, 1994.

- Dragan, Anca, Helen King, and Allan Dafoe, “Introducing the Frontier Safety Framework,” blog post, Google DeepMind, May 17, 2024.
- Financial Crisis Inquiry Commission, *The Financial Crisis Inquiry Report: Final Report of the National Commission on the Causes of the Financial and Economic Crisis in the United States*, U.S. Government Printing Office, 2011.
- Frelinger, David R., and Karl P. Mueller, *The AGI Rideout Strategy for Reducing Strategic Risk and Promoting Stability in the Transition to Artificial General Intelligence*, RAND Corporation, PE-A4347-1, April 2026. As of August 12, 2026:
<https://www.rand.org/pubs/perspectives/PEA4347-1.html>
- Future of Life Institute, “Pause Giant AI Experiments: An Open Letter,” March 22, 2023.
- Gaddis, John Lewis, *Strategies of Containment: A Critical Appraisal of American National Security Policy During the Cold War*, revised and expanded ed., Oxford University Press, 2005.
- Greenblatt, Ryan, Buck Shlegeris, Kshitij Sachan, and Fabien Roger, “AI Control: Improving Safety Despite Intentional Subversion,” arXiv, arXiv:2312.06942v5, July 23, 2024.
- Haasnoot, Marjolijn, Jan H. Kwakkel, Warren E. Walker, and Judith ter Maat, “Dynamic Adaptive Policy Pathways: A Method for Crafting Robust Decisions for a Deeply Uncertain World,” *Global Environmental Change*, Vol. 23, No. 2, April 2013.
- Hausenloy, Jason, Andrea Miotti, and Claire Dennis, “Multinational AGI Consortium (MAGIC): A Proposal for International Coordination on AI,” arXiv, arXiv:2310.09217, October 13, 2023.
- Hendrycks, Dan, Eric Schmidt, and Alexandr Wang, “Superintelligence Strategy: Expert Version,” arXiv, arXiv:2503.05628v2, April 14, 2025.
- Ho, Lewis, Joslyn Barnhart, Robert Trager, Yoshua Bengio, Miles Brundage, Allison Carnegie, Rumman Chowdhury, Allan Dafoe, Gillian Hadfield, Margaret Levi, Duncan Snidal, “International Institutions for Advanced AI,” arXiv, arXiv:2307.04699v2, July 11, 2023.
- Hoehn, Andrew, and Thom Shanker, *Age of Danger: Keeping America Safe in an Era of New Superpowers, New Weapons, and New Threats*, Hachette Books, 2023.
- Kapoor, Sayash, Rishi Bommasani, Kevin Klyman, Shayne Longpre, Ashwin Ramaswami, Peter Cihon, Aspen Hopkins, Kevin Bankston, Stella Biderman, Miranda Bogen, et al., “On the Societal Impact of Open Foundation Models,” arXiv, arXiv:2403.07918v1, February 27, 2024.
- Kokotajlo, Daniel, Scott Alexander, Thomas Larsen, Eli Lifland, and Romeo Dean, “AI 2027,” AI Futures Project, April 3, 2025.
- Korbak, Tomek, Joshua Clymer, Benjamin Hilton, Buck Shlegeris, and Geoffrey Irving, “A Sketch of an AI Control Safety Case,” arXiv, arXiv:2501.17315v1, January 28, 2025.
- Kulveit, Jan, Raymond Douglas, Nora Ammann, Deger Turan, David Krueger, and David Duvenaud, “Gradual Disempowerment: Systemic Existential Risks from Incremental AI Development,” arXiv, arXiv:2501.16946v2, January 29, 2025.
- Lempert, Robert J., Steven W. Popper, and Steven C. Bankes, *Shaping the Next One Hundred Years: New Methods for Quantitative, Long-Term Policy Analysis*, RAND Corporation, MR-1626-RPC, 2003. As of August 11, 2026:
https://www.rand.org/pubs/monograph_reports/MR1626.html

- Lindblom, Charles E., “The Science of ‘Muddling Through,’” *Public Administration Review*, Vol. 19, No. 2, Spring 1959.
- MacAskill, William, *What We Owe the Future*, Basic Books, 2022.
- MacAskill, William, “Viatopia,” *Forethought*, January 7, 2026.
- Marchau, Vincent A. W. J., Warren E. Walker, Pieter J. T. M. Bloemen, and Steven W. Popper, eds., *Decision Making Under Deep Uncertainty: From Theory to Practice*, Springer, 2019.
- Miotti, Andrea, Tolga Bilge, Dave Kasten, and James Newport, *A Narrow Path: How to Secure Our Future*, narrowpath.co, April 2025.
- Mitre, Jim, and Joel B. Predd, *Artificial General Intelligence’s Five Hard National Security Problems*, RAND Corporation, PE-A3691-4, February 2025. As of August 13, 2026: <https://www.rand.org/pubs/perspectives/PEA3691-4.html>
- Moon, Alvin, and Benjamin Boudreaux, *A Formal Model of How AI Erodes Human Agency*, RAND Corporation, RR-A4817-1, 2026. As of August 12, 2026: https://www.rand.org/pubs/research_reports/RRA4817-1.html
- OpenAI, “Preparedness Framework,” version 2, last updated April 15, 2025.
- Ord, Toby, *The Precipice: Existential Risk and the Future of Humanity*, Hachette, 2020.
- Predd, Joel B., James Baker, Benjamin Boudreaux, Edward Geist, and Matt Chesson, *Finding Common Ground in AGI Strategy Debates: A Diagnosis and Prescription*, RAND Corporation, PE-A4448-1, March 2026. As of August 12, 2026: <https://www.rand.org/pubs/perspectives/PEA4448-1.html>
- Predd, Joel B., Matt Chesson, and Gregory Smith, *Infinite Potential—Insights from the Two Moonshots Scenario: After-Action Report from a Sequence of Day After Artificial General Intelligence Exercises*, RAND Corporation, RR-A4230-1, 2025. As of August 12, 2026: https://www.rand.org/pubs/research_reports/RRA4230-1.html
- Rumelt, Richard, *Good Strategy/Bad Strategy: The Difference and Why It Matters*, Crown Business, 2011.
- Sajadieh, Sha, Loredana Fattorini, Raymond Perrault, Yolanda Gil, Vanessa Parli, Lapo Santarlasci, Juan Pava, Nestor Maslej, Russ Altman, Erik Brynjolfsson, et al., *Artificial Intelligence Index Report 2026*, AI Index Steering Committee, Institute for Human-Centered AI, Stanford University, April 2026.
- Sastry, Girish, Lennart Heim, Haydn Belfield, Markus Anderljung, Miles Brundage, Julian Hazell, Cullen O’Keefe, Gillian K. Hadfield, Richard Ngo, Konstantin Pilz, et al., “Computing Power and the Governance of Artificial Intelligence,” arXiv, arXiv:2402.08797v1, February 13, 2024.
- Schelling, Thomas C., *The Strategy of Conflict*, Harvard University Press, 1960.
- Schelling, Thomas C., *Arms and Influence*, Yale University Press, 1966.
- Sharma, Mrinank, Miles McCain, Raymond Douglas, and David Duvenaud, “Who’s in Charge? Disempowerment Patterns in Real-World LLM Usage,” arXiv, arXiv:2601.19062v1, January 27, 2026.
- Shavit, Yonadav, “What Does It Take to Catch a Chinchilla? Verifying Rules on Large-Scale Neural Network Training via Compute Monitoring,” arXiv, arXiv:2303.11341v2, May 30, 2023.

- Smith, Gregory, Benjamin Boudreaux, Michael J. D. Vermeer, and Joel B. Predd, *Infinite Potential—Insights from the Robot Insurgency Scenario: After-Action Report from a Sequence of Day After Artificial General Intelligence Exercises*, RAND Corporation, RR-A4231-1, 2025. As of July 30, 2026:
https://www.rand.org/pubs/research_reports/RRA4231-1.html
- Smith, Gregory, George Hage, Chad Heitzenrater, Matt Chesson, and Richard S. Girven, *Infinite Potential—Insights from the Cyber Surprise Scenario: Post-Series Scenario Report from a Sequence of Day After Artificial General Intelligence Exercises*, RAND Corporation, RR-A4626-1, 2026. As of July 30, 2026:
https://www.rand.org/pubs/research_reports/RRA4626-1.html
- Trager, Robert, Ben Harack, Anka Reuel, Allison Carnegie, Lennart Heim, Lewis Ho, Sarah Kreps, Ranjit Lall, Owen Larter, Seán Ó hÉigeartaigh, Simon Staffell, and José Jaime Villalobos, “International Governance of Civilian AI: A Jurisdictional Certification Approach,” arXiv, arXiv:2308.15514v2, September 11, 2023.
- Unruh, Gregory C., “Understanding Carbon Lock-In,” *Energy Policy*, Vol. 28, No. 12, October 1, 2000.
- U.S.–China Economic and Security Review Commission, *2024 Report to Congress*, November 2024.
- Verdon, Guillaume [@BasedBeffJezos], and @bayeslord, “Notes on e/acc Principles and Tenets,” *e/acc newsletter*, Substack, October 30, 2022.
- Vermeer, Michael J. D., *Historical Analogues That Can Inform AI Governance*, RAND Corporation, RR-A3408-1, 2024. As of August 12, 2026:
https://www.rand.org/pubs/research_reports/RRA3408-1.html
- White House, *Winning the Race: America’s AI Action Plan*, Executive Office of the President, July 23, 2025.
- Winter, Christoph, and Charlie Bullock, “Radical Optionality: Governing Transformative AI Under Uncertainty,” Institute for Law & AI, April 2026.
- Wohlstetter, Roberta, *Pearl Harbor: Warning and Decision*, Stanford University Press, 1962.
- Yudkowsky, Eliezer, and Nate Soares, *If Anyone Builds It, Everyone Dies: Why Superhuman AI Would Kill Us All*, Little, Brown and Company, 2025.

About the Authors

Joel B. Predd is a senior engineer at RAND and director of the Center for the Geopolitics of AGI. His research focuses on the geopolitics of AI and on the intersection of the security and economic competition with China. He holds a Ph.D. in electrical engineering.

Benjamin Boudreaux is a policy researcher at RAND and professor at the RAND School of Public Policy. His research focuses on the ethics of emerging technology. He holds a Ph.D. in philosophy.

Matt Chessen is an expert technical resident at RAND. His research focuses on the geopolitics of AGI, artificial superintelligence, and scenario development and planning, among other topics. He holds a J.D. and an M.B.A.

Beba Cibralic is a policy researcher at RAND and a professor of policy analysis at the RAND School of Public Policy. Her research focuses on AI governance, ethics, and policy. She holds a Ph.D. in philosophy.

Edward Geist is a senior policy researcher at RAND. His research interests include the former Soviet Union, nuclear weapons, emergency management, and AI. He holds a Ph.D. in Russian history.

Kamaria Horton is a technical analyst at RAND. Her research focuses on AI and its intersection with governance, geopolitics, security, and safety. She holds an M.A. in global affairs.

Amanda Kerrigan is a policy researcher at RAND. Her research focuses on social stability in China, development of AI in China, and other China-related topics. She holds a Ph.D. in China studies.

William Marcellino is a senior behavioral scientist at RAND. He develops natural language and AI tools and systems and conducts policy research on such topics as AI, foreign malign information operations, and military information technology systems. He holds a Ph.D. in rhetoric.

Shanshan Mei is a political scientist at RAND. Her research focuses on Indo-Pacific Security and the People's Liberation Army. She holds a Ph.D. in international relations.

Jared Mondschein is a physical scientist at RAND. His research focuses on critical and emerging technologies, AI diplomacy, and supply chains. He holds a Ph.D. in chemistry.

Alvin Moon is a mathematician at RAND. His research focuses on mathematical analysis of AI, emerging technologies, cryptography, and supply chains, among other areas. He holds a Ph.D. in mathematics.

Tobias Sytsma is an economist at RAND. He conducts research on topics around technological change, international trade, and economic security. Sytsma holds a Ph.D. in economics.