

Issue Brief

Artificial Intelligence in Competition

How Will AI Change the Game and
What Does It Mean to “Win” in AI?

Author

Andrew J. Lohn



CSET CENTER *for* SECURITY *and*
EMERGING TECHNOLOGY

September 2026

Executive Summary

Policymakers are trying to win a race for AI, but neither developing nor deploying technology is a sufficient goal on its own. This brief contends that the goal should be to achieve grander victories such as in economic or military domains, with AI as an enabling technology. Achieving that larger goal demands success in two additional stages beyond AI development and deployment. First, AI needs to provide some specific competitive advantage. After that, policymakers must develop a plan for, or theory of, what constitutes victory, describing how competitive advantages lead to triumph despite an adversary's efforts to the contrary.

Figure E.1: The Chain From AI Superiority to Winnings Can Break in Many Ways



Source: Author's analysis.

AI can provide many different types of competitive advantages to varying degrees, depending on the application. In some cases, AI will be decisive. In others, it will provide advantages that are only substantial, negligible, or even counterproductive. There are too many possible applications of AI to evaluate each explicitly, so eight different types of competitive advantages are considered here, as listed in Figure E.2. For each type, illustrative examples are provided from game theory, athletics, chess, cybersecurity, or finance, to clarify how an AI advantage can be decisive, substantial, negligible, or counterproductive.

Figure E.2: Eight Different Types of Competitive Advantage Are Considered

<u>Information Asymmetry</u> 	<u>First Mover</u> 	<u>Second Mover</u> 	<u>More Options</u>
<u>Fewer Options</u> 	<u>More Time</u> 	<u>More Materiel</u> 	<u>Differing Objectives</u>

Source: Author's analysis.

Based on those examples, AI is likely to provide some decisive advantages, but the number of decisive cases will be limited by non-AI limitations, bottlenecks, and adversary countermeasures. Additionally, the degree of AI superiority may be too small and fleeting to provide many decisive advantages.

Even a decisive competitive advantage may not be enough to secure victory. Achieving strategic-level goals such as in military conflict, diplomacy, or profit and loss requires another level of prudence. For any specific AI application, there are six categories of theory of victory to choose from.

- **Cooperation** looks for win-wins, where a rising tide lifts all boats. The supply chain could be spread across partners, and the rewards from AI could be widely distributed rather than fought for.
- **Dominance** exists when a stronger competitor can enforce its will on a weaker competitor with little recourse. Dominance has been the United States' primary approach to AI at both the geopolitical and corporate levels.
- **Denial** prevents adversaries from achieving their goals. Denial has also been a popular theory of victory for the United States in AI through export controls and protection of intellectual property such as model weights.
- **Devaluing** reduces the benefits of victory in ways similar to burning crops in retreat so that an advancing army gains little. In AI, a devaluer might reduce markets where AI would succeed, or degrade an information environment so that AI becomes less useful.
- **Brinkmanship** pushes a risk of mutual harm to the point where an adversary relinquishes. It has led to arms races in the past that could also play out in AI, for example to develop weapon stockpiles or potentially dangerous AI capabilities.
- **Cost Imposition** drives rivals to spend more than they are capable of sustaining. It was part of a theory of victory during the Cold War and is increasingly pertinent in AI, as expenses balloon into the trillions of dollars.

These types of theories of victory are general enough to apply to any strategic goal. This report applies them to assess which approaches the United States can take to achieve outsized rewards from AI. The report also anticipates approaches that adversaries might take to reduce those rewards or achieve their own goals.

- **Cooperation** does not currently look promising among rivals amid geopolitical decoupling and increasingly secretive companies. On the deployment side, there are also reasons to be suspicious of cooperation, including bottlenecks that could prevent AI from “increasing the pie.” Still, progress draws heavily on open science and contributions from geopolitical allies such as the Netherlands, Taiwan, Japan, and South Korea, among others.
- **Dominance** is challenging, given that technological leads are small and shrinking. The set of cases inspected in this report also suggests that there will be relatively few durable and decisive competitive advantages from AI.
- **Denial**’s export controls and increased secrecy have not prevented weaker competitors from making progress. From the opposite perspective, weaker competitors may adopt a denial strategy by deploying AI that may not be good enough to win outright, but is good enough to deny victory to stronger rivals.
- **Devaluing** may also be a strategy for weaker competitors who could degrade markets for AI by releasing good-enough alternatives. Weaker competitors can also act chaotically or poison AI systems to reduce disadvantages.
- **Brinkmanship** does not appear viable because the anticipated benefits are too high and the risks too dispersed. As nations build stockpiles of AI-enabled drones and companies race to develop systems that they say might destroy humanity, nobody seems to be backing down.
- **Cost Imposition** is an increasingly plausible theory of victory for weaker competitors because diminishing returns force AI leaders to spend more to maintain smaller leads in both development and deployment.

It is not enough to win the AI race to develop and deploy AI. AI must provide desirable outcomes in the presence of adaptive adversaries. It is not an easy task to ensure desirable outcomes because although there are many ways that AI might provide an advantage, there are also many ways for an advantage to be smaller than anticipated, or for adversaries to intervene. Policymakers need a step-by-step process for identifying these competitive advantages, disadvantages, and adversary actions. This brief aspires to provide that process, converting AI superiority into theories of victory that government or corporate leaders can use to achieve AI’s promise.

Table of Contents

Executive Summary.....	1
Introduction.....	5
What Does It Mean to Win?.....	8
Changing the Game	11
Types of Advantage.....	13
Theories of Victory.....	17
Conclusion.....	22
Author.....	23
Acknowledgments.....	23
Appendix A: Competitive Advantages.....	24
Appendix B: Walkthrough of AI Superiority to Victory in Construction Materials Industry	41
Appendix C: Walkthrough of AI Superiority to Victory in Conflict Between Large and Medium-Sized Militaries.....	47
Endnotes.....	57

Introduction

At the strategic level in the United States, AI is largely viewed as a race to be won. The AI Action Plan is titled “Winning the Race.”¹ But in AI, “winning” is a complex goal. Victory should not be defined in terms of AI itself but rather in terms of the benefit, such as economic profits, military objectives, or diplomatic gains.

In addition to the U.S. government, companies and foreign countries are racing to develop the best AI systems, but few can explain how being first in the race for AI translates into profit, winnings, or net benefit. Companies that have been quick to deploy AI have often been surprised by failed pilot projects and initial decreases in productivity.² Whether the competition is economic, military, political, or social in nature, and whether it is among nations, companies, or individuals, the path from AI progress to successful outcomes is less direct than it might appear.

The Chain from Superiority to Victory

This report breaks the chain that links AI progress to victory into three stages. The first is achieving AI superiority. This is the often-discussed race to develop and deploy AI. The term “AI race” refers to many now-familiar aspects of the race, such as financial investments, chip production, the role of open source, energy and water supply, and many others. But achieving superiority in AI development and dissemination is often where the discussion ends.

Figure 1: Achieving Victory After Winning the AI Race



Source: Author's analysis.

The second stage in the chain is turning AI superiority into a competitive advantage. This report evaluates eight different types of competitive advantage, as listed in Figure 2, and discusses AI's role in providing, increasing, or reducing those advantages. The final stage is a plan or process for turning competitive advantage into victory.

A theory of victory lays out a causal chain of actions and countermeasures describing how competitors use their advantages and opportunities to achieve their goals. All causal chains for all competitions cannot be listed here. Instead, this report considers

the six different categories of theory of victory listed in Figure 3, and describes how the competitive advantages from the second stage of the chain enable or detract from each theory of victory.

Figure 2: Eight Different Types of Competitive Advantage Are Considered

<u>Information Asymmetry</u> 	<u>First Mover</u> 	<u>Second Mover</u> 	<u>More Options</u> 
<u>Fewer Options</u> 	<u>More Time</u> 	<u>More Material</u> 	<u>Differing Objectives</u> 

Source: Author’s analysis.

Figure 3: Six Categories of Theory of Victory Are Considered

<u>Cooperation</u> 	<u>Dominance</u> 	<u>Denial</u> 
<u>Devaluing</u> 	<u>Brinkmanship</u> 	<u>Cost Imposition</u> 

Source: Author’s analysis.

Failure is possible at each stage. In the first stage, the United States could fail to achieve AI superiority in any of its various dimensions. American AI offerings may not continue to lead in capability, or they could fall behind in other ways such as speed, cost, reliability, size of user base, or embeddedness throughout the global ecosystem. If AI superiority is achieved, it may then fail to provide competitive advantage in the second stage. AI may be less useful than some anticipate and may run into limiting bottlenecks or barriers in deployment and use. Or it may be too costly to be profitable.³ It may force adversaries to take undesirable countermeasures such as degrading the information environment or developing more autonomous and rapid countermeasures. And AI may simply add little when existing solutions are already sufficient.

Achieving the third stage requires converting competitive advantages into victory. Even overwhelming competitive advantages, whether economic, military, or technological, do not always translate into strategic wins, as the United States learned in Vietnam and the Soviet Union learned in Afghanistan. Military strategists now look back and question why decades of overwhelming military capabilities provided few strategic victories. Technology strategists should look forward, accounting for likely adversary countermeasures and asking how AI advantages will produce strategic victories.

Structure of the Report

Before describing how AI affects winnings, the next section of this report briefly reviews what it means to win. Following that, ways in which AI can be literally game-changing are explored, illustrating that it is possible for progress to be detrimental for all parties. Building on AI's potential to be beneficial or detrimental, eight types of competitive advantages are described, including how AI can make those advantages decisive, substantial, negligible, or counterproductive. The 32 (8 advantages times 4 degrees of advantage) are detailed in Appendix A. Next, the viability of six different theories of victory is considered, incorporating the analysis of competitive advantages. And finally, this report's conclusions are summarized.

The body of the report focuses primarily on AI competition at the overall national level to assist federal or state-level decisionmakers. The approach also applies to more narrow competitions. Appendix B provides an example of the chain from AI superiority to competitive advantage to theory of victory in a construction-materials industry such as steel or concrete, and Appendix C does the same for a conflict between large and medium-sized militaries.

What Does It Mean to Win?

Philosophers have grappled with what it means to win for thousands of years, since at least 280 BCE, when King Pyrrhus of Epirus is said to have exclaimed, “If we are victorious in one more battle with the Romans, we shall be utterly ruined.”⁴ More recently, winning and losing has been a focus for game theorists, economists, international-relations scholars, evolutionary biologists, and many others.⁵ Because AI can be applied so widely, victory in this domain can mean many things.

The literature review for this report did not turn up a definitive summary of what it means to win across the many fields. What follows is not comprehensive, it only aims to concisely refamiliarize readers with a range of interpretations of “winning.”

Win-Wins and Zero Sum

If winning simply means receiving more benefit than cost, then multiple actors can succeed, enabling cooperation and win-wins. For example, AI could increase the size of a competitive market. If one company captures most of that market, another can still win by increasing revenue despite losing market share.

Alternatively, winning can mean receiving the most net benefit. In that case, a competitor could win by racing ahead or by suppressing opponents, possibly even if that comes with negative net benefit to itself. Two countries at war both suffer casualties, the winner may be the one who suffers least from the losses.

Increasing Payoffs

In any game, each possible outcome has an associated payoff that could be in the form of money, happiness, or territory, for example. Sometimes the payoffs are discrete, such as in chess, where a player can only win, lose, or draw. In other cases, the payoff may vary more continuously, such as in business with net profits, or in conflict with territory gained. In either case, AI can be envisioned to increase payoffs. It may increase revenue, decrease cost, or, as in chess, increase the odds of discrete wins.

More Iterations

AI can also increase winnings without increasing payoffs if it can increase the number of times that a beneficial game is played. If an AI system reduces the time per iteration, then the game can be played more times. The result is more winnings, potentially for everyone involved. Accelerating financial services or bureaucratic processes may have this effect.

There are many ways in which AI could help games be played more times. For example, with a limited amount of computing power, a more efficient AI system could iterate a competition more times. Similarly, if supervisor time is a limiting factor, then as an AI system becomes more reliable, the game can be iterated more times for a given number of human supervisors. But there can also be trade-offs, such as between speed or efficiency and between quality or reliability.⁶

Achieving the Value of the Game

Each iteration of a game has a “value,” which is the amount that should be expected for that iteration if the game is played perfectly, for example how many dollars the dealer in a heads-up game of poker should expect to win per hand.⁷ AI may help competitors play more perfectly to achieve the value of the game, thereby increasing winnings. But once competitors can achieve the value of the game, AI does not help. Many games, such as checkers, tic-tac-toe, and to a lesser degree poker, are “solved” in the sense that the ideal moves are already understood, providing little opportunity for further improvement.⁸

At the other extreme, many contests are so chaotic, complex, or random that AI is unlikely to help a competitor achieve the value of the game, even if a perfect model is developed.⁹ Some military battles may depend on too many important and unknowable variables for even a powerful AI model to be much help. And a weaker competitor may try to level the playing field by introducing additional randomness or chaotic elements about which neither player can reason. A weaker competitor may even make military or corporate decisions on a dice roll that AI cannot help anticipate. Or businesses and governments could release troves of fake documents to muddy the information environment. These types of behaviors can nullify an AI advantage, but can also be mutually detrimental.

Outwitting to Surpass the Value of the Game

The value of a game is not always the upper bound on how much a competitor can expect to win. If a game is not solved, then a weaker competitor may make mistakes. Armed with a superior AI, competitors might tempt weaker opponents into mistakes, or capitalize on blunders. However, those gambits can be risky if tempting an opponent to make errors creates opportunities for that opponent, or if what appears to be a blunder may instead be a trap. AI systems are notoriously vulnerable to manipulation, so traps such as military deception, subterfuge such as data poisoning, or manipulation such as prompt injection may be common even against highly capable AI systems.¹⁰ Powerful

defensive cyber agents may still be easy for attackers to manipulate into releasing sensitive information or disabling critical systems.

Cooperation to Surpass the Value of the Game

Human competitors often choose “irrational” moves that weaken their position for good reasons. Humans act altruistically or punish antisocial behavior in ways that might seem irrational outside of the broader context that includes factors such as the value of reputation or the satisfaction of being a good person.

The Dictator game is a widely studied example.¹¹ In it, one player splits a pot between himself and other participants. The mathematically optimal choice is for the dictator to keep the entire pot, but humans often altruistically offer some payoff to other players. The Ultimatum game proceeds the same way, except that the second player can either accept or reject the offer.¹² If the offer is rejected, then neither player receives anything. The optimal play within the context of this game is to accept even small offers, but humans often reject offers that they feel are unfair in order to punish the other player. These types of dynamics can play out in negotiations and agreements large and small that humans conduct today but that AI might manage in the future.

A slightly more complex scenario is that of the Centipede game, where the pot grows over time.¹³ The first player to split the pot receives a disproportionate share, but everyone benefits from waiting before splitting the pot. The game’s optimal move is to split the pot at the first opportunity, knowing that the next player is always incentivized to do so. Instead, humans often grow the pot for several iterations, to everyone’s advantage.¹⁴ This scenario can play out in practice, such as in an open-source software project where the value of the project increases with each contribution but the majority of the profit goes to the first contributor to commercialize it.

AI systems might tend toward the mathematically optimal within-game decision more often than the socially beneficial option that humans prefer. On the other hand, AI systems might retain some of these socially beneficial irrationalities. AI that has been trained on human data does incorporate some of these behaviors and can be prompted to act more or less socially.¹⁵ A related concern, though, is whether humans will retain their socially beneficial irrationalities as AI becomes more prevalent. In general, humans act less altruistically toward machines, but that too can change depending on the human’s mindset or perceptions of the machine.¹⁶

Changing the Game

AI may increase payoffs in ways that literally change the game, and not necessarily for anyone's benefit. This can be illustrated in simple two-player games where each player can only cooperate or defect. It is a popular structure in game theory because despite its simplicity, it helps demonstrate or explain behaviors that are common in more complicated real-world competitions.

The game can be visualized using payoff tables such as the one in Table 1. Both players receive a payoff of a if they both cooperate, and d if they both defect. If one player defects and the other attempts to cooperate, the defector receives b and the cooperator receives c . Depending on the relative values of a , b , c , and d , this structure becomes different well-known games such as the Prisoner's Dilemma, Stag Hunt, Chicken, and a version of Battle of the Sexes.

Table 1: Payoffs for a Common Structure in Two-Player Games

		Player 2	
		Cooperate	Defect
Player 1	Cooperate	a,a	c,b
	Defect	b,c	d,d

Source: Author's analysis.

In the case where those payoffs are ordered as $b > a > d > c$, the game is called a Prisoner's Dilemma. In its description, two suspects are arrested and offered plea bargains to implicate each other. With those payoffs, their best choice would be to cooperate and stay silent, but they are both incentivized to defect, leading to worse outcomes for both.

If instead the payoffs are ordered as $a > b > d > c$, the game is called a Stag Hunt. Two hunters can cooperate to catch a deer, or they can hunt alone to catch rabbits, but they do poorly if they hunt deer alone. Game theoretically, the incentives to defect still

remain, but the Stag Hunt has a second viable strategy. The incentives also allow competitors to cooperate, resulting in the largest possible benefits for both.

These two games have drastically different strategies and outcomes but are nearly identical in structure. Changes in payoffs can convert one to another. The Stag Hunt is clearly preferable, and the Prisoner's Dilemma is to be avoided. Worryingly, small *increases* in payoffs can convert a Stag Hunt that incentivizes cooperation toward mutually maximizing benefits into a Prisoner's Dilemma incentivizing defection with mutually punishing outcomes. That switch from cooperation to defection happens if the payoff for going it alone (*b*) rises above the payoff for cooperating (*a*). The crossover point is a threshold that can lead to collapse.

As a tangible AI-relevant example, imagine two news companies that can either use human reporters to create high-value content or AI bots that produce lower-quality content at lower cost. Further, imagine that if either company uses exclusively AI bots, then the cost pressure to compete will make human-reported news unsustainable. This scenario is a Stag Hunt that promotes collaboration. If one company chooses human generation, then the other should too. But the game flips to the defection-dominant Prisoner's Dilemma if the payoff for any company using AI while others use humans becomes greater than for everyone using humans. That could happen if AI is inexpensive but trains on the outputs of the other companies. Those cost savings incentivize all companies to use AI-only generation, even if that means none will be able to sustain human-generated news, and all will earn less as a result of the degraded training data.

Similar scenarios can be envisioned for all sorts of markets, including art, music, and science. It may be a common mechanism by which advances in AI could lead to lower profit for all companies involved, and worse outcomes for customers. And it may be playing out across the internet as a whole, as patrons turn to AI outputs rather than visiting and supporting the webpages that produce the original information on which AI relies.¹⁷

Types of Advantage

Competitors strive to win by developing advantages that can come in many different forms. Eight different types of advantages are considered here. For each type, this report also considers how AI can make those advantages decisive, substantial, negligible, or counterproductive. The result of this analysis is 32 different advantage-degree combinations, detailed in Appendix A. This section provides a brief summary of each type of advantage. Readers are encouraged to refer to the appendix for a more complete description of the examples and lines of reasoning for the following summaries.

Information Asymmetry

An information asymmetry may exist if a competitor has access to more information, or if the competitor can make better sense of that information. AI is built to analyze and generate information, so it is well suited to these information tasks but also faces limitations. These AI-enabled information asymmetries can be decisive if they reveal few-point failures that can be acted on, revealing some secret information or vulnerability that shifts the tides and would not have been revealed otherwise. More often, information advantage is not so critical. That may be because it cannot be acted on, which is why nations are sometimes comfortable telegraphing their missile strikes in conflict. Alternatively, the information advantage may be fragile or short-lived, for example how a competitor can change encryption schemes after a communications leak is revealed. In theory, information advantages can also be counterproductive, including by disincentivizing negotiation in military crises when one nation is certain it will win as long as it does not reveal the information that gives the certainty.¹⁸ AI can also degrade trust between groups, for instance, as famously studied, in used car sales where well-informed salesmen become unwilling to part with what they know are quality vehicles at prices that less-knowledgeable customers are willing to accept, resulting in market degradation.¹⁹

First Mover Advantage

AI promises superhuman speeds in the ability not only to act autonomously, but also to accelerate human decision-making and business processes. It may provide the speed to win races, or the initiative to lock in a lead such as with patents, or to compound early benefits such as attracting more customer data or revenue to reinvest. But the rush to move first can also exacerbate safety or reliability risks.²⁰ It inflates the financial and environmental costs of AI development while increasing time pressures that impede crisis stability, oversight, and de-escalation off-ramps.²¹ First mover incentives

might also drive AI, or humans in the presence of AI, to logically choose mutually worse outcomes in collaborations, as in the Centipede game discussed earlier. On the other hand, advances in AI can reduce first mover advantages by improving responsiveness. As a simple and tangible example of degrading first mover advantages, in chess AI appears to be helping black play to draws.²²

Second Mover Advantage

Second mover advantages may be less apparent than first mover advantages, but they are common nonetheless. AI may help second movers find flaws in physical defenses such as the infamous Maginot line, or in digital defenses in cybersecurity.²³ AI may also provide the responsiveness to win in Rock-Paper-Scissors, a game of Chicken, or any situation where the best move becomes obvious once the first is played.²⁴ In AI development itself, or technology deployment in general, second movers can cut costs and need to innovate less, benefiting from others' investments and the trend of falling prices. But AI could also weaken second mover advantages by improving initial bids, for example, leaving little on which a second mover could capitalize. And AI could increase the odds of deadlock by incentivizing all competitors to wait.

More Options

Borrowing from chess parlance, AI may turn Pawns into Queens, increasing the versatility and capability of employees, soldiers, or software. In chess that would be enough to nearly secure victory, but more often competitors will be able to adapt to limit those advantages and exploit the weaknesses in AI systems. For example, adding an additional and more capable move to Rock-Paper-Scissors that defeats two moves instead of the usual one provides little value. That is because rivals can effectively adapt to limit the impact of that new move even in such a constrained contest.²⁵ In the push to modernize and cash in on new highly capable technology, the benefits of AI may not always be as large as they seem after accounting for countermeasures. Given AI's notorious vulnerability to adversarial manipulation, even very weak rivals will typically have at least some countermeasures to which they can turn.

Fewer Options

Perhaps counterintuitively, having fewer options to choose from can also be an advantage.²⁶ As humans delegate more decisions to AI agents, those agents will be constrained in various ways. Whether booking hotel rooms on a capped budget or fighting to the robot death without approval to retreat, a more limited AI system can use those limits as negotiating leverage to its advantage. In stark cases, such as *Dr.*

Strangelove's Doomsday Device, the advantage could be decisive, but the risks of brinkmanship are typically not worth the benefits. More generally, AI increases adaptability rather than restricts it, so AI may not provide a fewer-options advantage often.

More Time per Turn

Much of this section can be related to others. Having more time, or having the speed to make better use of time, can help win races, as described in the section about first mover advantages. More time can also provide more responsiveness or adaptability, as described in the section about second mover advantages. More time can provide more options if it allows for more actions per minute or more actions before a competitor can respond. Time can also provide more options by accelerating aspects of the process of developing options, for example identifying alternatives, soliciting feedback, or anticipating outcomes. The advantages of time that can be cast to the other advantages have already been discussed. But other ways in which time can be advantageous have not yet been covered.

AI might accelerate the stages of the Observe-Orient-Decide-Act (OODA) loop for more than just fighter pilots.²⁷ It might also give competitors the precision to make last-minute decisions, leaving little time for rivals to act. On the other hand, if AI helps make decisions quickly, then a short-clocked adversary may find the remaining time perfectly sufficient. Further, decision speed is not a limiting factor when other delays are included, such as internet bandwidth, human review in negotiation, or material shipments in logistics or missile flight times. Even where time is a primary factor, rather than consume time, AI may provide an impetus for instituting new time controls, as it has already done for throttling download speeds and implementing stock market circuit breakers. Carefully constructed, those time controls can help incentivize good decisions rather than forcing rash ones.

More Material/Materiel

Game theory's Colonel Blotto decides on quantitative trade-offs in military scenarios, election campaigning, and many other domains.²⁸ In the game, Colonel Blotto has a total number of armies, drones, or dollars to allocate across various battlefields without knowing how much the adversary will allocate to those battles. A player wins battles in which the allocation exceeds the enemy's, and loses where it is lower.

Colonel Blotto's dilemma would be solved if AI provides enough military drones or economic production to overwhelm an adversary. But superior numbers do not

necessarily provide superior might. It may be that a military's defensive three-to-one advantage, or a cyber defenders' layered barriers, set too high a bar for an AI-enabled material advantage to be decisive. Gene Amdahl's observations about bottlenecks and Robert Solow's about limited economic productivity gains might continue to hold true for AI in ways that result in less growth than many anticipate.²⁹ In that case, the costs of AI development may be regrettable.

Differing Objectives

The advantages discussed so far implicitly assume that competitors are trying to achieve similar goals. An entirely different advantage comes from having an easier objective. For example, it may be easier to hide than to seek. In games where the competitors have different goals, the environment often matters more than the players' skill, and AI is likely to change both.

Advances in AI can help find the needle in some haystacks, such as faces in a crowd, but in many cases, it is easy to increase the size of the haystack or adapt in other ways to compensate. That has been true throughout technology transitions, as competitors in cyber and trench warfare have adapted to maintain balance. In some cases, though, the adaptations can exacerbate risks. Nuclear-armed states that move to a hair trigger in fear that AI will locate their weapons increase the danger for everyone.³⁰ And incentivizing deception or uncertainty to throw off AI systems risks mistakes and miscalculations.³¹

Theories of Victory

It is not enough to understand AI's advantages and what it means to win. Without a strategy and theory of victory, AI may provide tactical victories and strategic defeat. Whether the objective is to increase net profits or to have more military might than an adversary, victory will occur through a causal chain linking the technology to eventual victory. That chain must include how the advantages will be used, how weaknesses will be mitigated, and how competitors or adversaries will respond.³²

RAND Corporation surveyed the literature to identify five categories of theories of victory.³³ The categories were developed in the context of wartime operations but can be viewed as “universal to all decisionmakers and to any conflict.” They are: Dominance, Denial, Devaluing, Brinkmanship, and Cost Imposition. Expanding beyond the wartime context, this report adds Cooperation as an additional category, highlighting that conflict is not inevitable and that there could be ways to win without others losing.

Cooperation

As companies and countries act to develop AI, they could draw on partnerships among one another. AI deployment too could be widespread, leading to widely distributed increases in productivity and prosperity rather than driving further competition for slices of a still-limited pie. And ideally, AI advances would benefit defense more than offense, decreasing the incentives for antagonism.³⁴ But there are many hurdles to cooperation that will likely require proactive steps from policymakers, developers, deployers, and users.

As AI increases payoffs, it must not change the game in ways that incentivize defection over cooperation. AI can turn cooperative Stag Hunts to Prisoner's Dilemmas by degrading the market for data in specific domains such as news media, scientific fields, art, and across the internet in general.

Beyond economic incentives, policymakers, developers, deployers, and users must protect the social incentives for cooperation. Whether it is a Centipede game such as collective development of open-source software or a Dictator game of distributing rewards, cooperation relies on behaviors that reinforce the social good in ways that machines may be less likely to choose.

This important theory of victory is fragile, but there are signs of promise. Open science still contributes significantly to AI development.³⁵ Open models are competitive with the less cooperative private models. And despite tensions and challenges among

geopolitical rivals, AI development relies heavily on cooperation among allies. The United States designs the best AI chips—that use semiconductors fabricated in Taiwan, that are patterned using machines built in the Netherlands, and that rely on chemicals from Japan and South Korea.

Decades of “normalization” of relations with China should be a cautionary tale about the downsides of being overly optimistic about cooperation.³⁶ Still, there are many opportunities for collaboration among allies and even some among rivals.

Dominance

Dominance as a theory of victory involves overwhelming the adversary. It is based on the presumption that the dominator’s weaknesses are insignificant and that the adversary’s countermeasures are insufficient. There are cases where those are likely to be true. Some examples are races, overwhelming materiel advantages, locked-in early leads such as from patents or having a chess board full of Queens face off against Pawns. But these dominant cases may be rarer than they may seem.

Real-world competitions tend to be complex, and the adversary has many opportunities to adjust its strategy or the rules of the game. Competitors can implement bandwidth limitations for cybersecurity or stock market circuit breakers to avoid races. They can lean into adversarial manipulations of vulnerable AI systems, reducing AI’s advantage similarly to how playing more Scissors limits the impact of adding Dynamite to Rock-Paper-Scissors. And weaker competitors can increase uncertainty, randomness, deception, or chaos to minimize advantages that AI offers.

Even vastly superior AI may not be decisive. AI is useful in competitions that are neither too easy nor too difficult. AI does not help in games that are already solved, such as checkers or poker. It also may not help for excessively complex or impossible tasks such as those in Clausewitz’s chaotic vision of war that is “everywhere brought into contact with chance, and thus incidents take place upon which it was impossible to calculate.”³⁷ The sweet spot for AI is chess that falls between checkers and Clausewitz. That sweet spot may be small, and even where it exists, the degree of AI superiority may not be enough to support dominance as a theory of victory. Even weak adversaries will have access to capable AI.

Further, even when AI advantages are decisive, strategic dominance can be difficult. Disadvantaged adversaries will adjust to minimize the importance of their disadvantage. They will shift emphasis away from battles where they can be overwhelmed by superior numbers. They will shift competitions toward bottlenecks, perhaps targeting supply chains or forcing superior adversaries to take serial actions,

one step after the other—approaches taken by the Viet Cong during the Vietnam War and more recently by the Taliban in Afghanistan.³⁸ Just as decades of overwhelming military might may have led to few strategic successes, even overwhelming AI advantages do not make a viable theory of victory on their own.³⁹

Denial

Denial is victory in the sense that the opponent does not win. It is more a theory of defeat than a theory of victory. In the context of AI, it could involve denying: (1) access to AI systems and the technology needed to develop or deploy them; (2) an advantage from those AI systems; (3) the rewards that come from those advantages.

Policy and debate in the United States often focus on denying access to AI through restrictions on computer chips, the models themselves, or the datasets and human talent to develop them. The U.S. government and some leading companies have made substantial efforts to protect intellectual property, corral talent, and restrict international transfers, but competitors have not been easily denied.⁴⁰ The current gap that separates closely guarded state-of-the-art AI systems from freely available alternatives or international competitors is not large enough to be confident in denial as a theory of victory.

A weaker competitor can focus on developing “good enough” AI. As in the example of “ ϵ -solving” poker, the possible winnings of even perfect play may be reduced to just the small number ϵ . In chess too, a good enough AI may not win many games, but it might avoid errors and play to a draw, still denying victory. It is unclear how many important competitions play toward a stalemate, but it appears that many AI capabilities face diminishing returns.⁴¹ Where there are diminishing returns, a good enough AI can deny a superior AI system from continuing to achieve much value as the cost to stay ahead grows.⁴²

Critically, stronger competitors must not let the tactical pursuit of developing AI blind them to whatever strategic objective they have for it because a weaker competitor can deny at the strategic level as well. An inferior AI can easily return more profit if it is less costly.⁴³ Or rather than profit, the objective may be influence, where wide dissemination of AI systems is more important than performance, cost, or revenue.⁴⁴ As the world becomes more reliant on AI, it may be the most deeply entrenched models and systems, rather than the most capable, that provide the most leverage in international negotiations or price-setting. Weaker competitors, forced to grapple with their capability limits, easily recognize that it may be the cheapest, fastest, most accessible, or least resource-intensive systems and models that have the largest effect

on the world. They will use those opportunities to deny the victory that stronger competitors anticipate.

Devaluing

Devaluing is also a theory of defeat more than a theory of victory and is likely to involve decreased benefits for both parties. A retreating army's scorched-earth approach to devaluing leaves nothing behind for the encroaching enemy to gain.

At the time of writing, many observers are reducing their expectations of AI's near-term value. That is mostly the result of rising costs and diminishing returns.⁴⁵ It is also the result of some of the pressures of Solow's paradox, where digital transformations can be slow to materialize or can devalue the markets on which they are meant to capitalize. But AI's perceived decrease in value may also be partly the result of more complex strategic interplay.

Rather than allow a competitor to win in a market or battle, a weaker competitor may degrade that market or destroy the territory they would lose. An effect of Chinese open-weights models is that they degrade the market for AI systems.⁴⁶ Developing a frontier AI system becomes a Pyrrhic victory if a weaker adversary releases a good enough alternative for free—as happened to the superior Netscape browser during the 1990s browser war, when Microsoft bundled and freely distributed its Explorer browser. Stronger competitors such as those in the United States need to recognize where they are vulnerable to devaluing in the technologies they develop or in the markets to which they deploy.

Brinkmanship

A competitor can force its adversary to accede by threatening punitive action that may even be self-destructive. AI might enable brinkmanship as a theory of victory in many ways including arms racing to develop AI-enabled systems such as drones, or racing to develop AI models themselves. Brinkmanship could also involve deploying agents or systems with high risk tolerance or escalatory guidance. It could involve negotiating with caps where a bargaining agent might withdraw rather than compromise. Or it could involve consuming time, as in games of Chicken.

AI also enables countermeasures to brinkmanship. It can reduce response times and increase adaptability, making it harder to constrain an adversary to an unacceptable outcome that would force it to submit. If AI can increase the number of options that are available, then it may decrease the need for, or the ability to attempt, brinkmanship.

Cost Imposition

Rather than reducing value, a victor can cripple an adversary with expenses or convince them that a prospective path would be too costly to pursue.⁴⁷ This was a strategy used by the United States against the Soviet Union in the 1980s.⁴⁸ AI expenses have ballooned to reach strategically relevant numbers. Companies in the United States are spending nearly \$1 trillion to develop and deploy models that are only marginally and fleetingly superior to the much less costly alternatives that fast-followers provide.⁴⁹ By continually releasing comparable models for free use that are also inexpensive to produce, China may be pursuing cost imposition as a theory of victory in AI against the United States.

Beyond the development of AI, deployment can be costly. AI may lead to more battles of attrition, forcing Colonel Blotto to spend more than he would prefer on AI-enabled drones, marketing dollars to run AI-written target ads, or datacenters to run cyberattack and defense agents. Policymakers need to be cognizant that the pressure they feel to outspend rivals may be part of an adversary's theory of victory.

Conclusion

AI promises to provide immense benefits, but those promises are not guarantees. There are many ways for AI to fall short of its potential, or for its potential to be overstated. Policymakers at every level, from proprietors of small companies to national leaders, must make many decisions about complex technology and even more complex sociological and competitive dynamics. Having and deploying the best AI systems is not enough.

To achieve victories that are worth the effort and expense, decisionmakers need to assess AI's impact through several stages, from winning the AI race to achieving victory. They need to identify which competitive advantages a superior AI will provide in a given context, justify whether those advantages will be substantial or decisive, and anticipate how those advantages will lead to victory in the presence of adaptive adversaries. This is a lot to think through, much of which can be counterintuitive. In some cases, the line of reasoning will be trivial, and AI will clearly lead to benefits. Many other cases will require meticulous planning and a step-by-step guide such as the one provided in this brief, helping walk through the stages from AI superiority to theory of victory. Such a guide should outline the main competitive advantages to consider, frame the types of counterintuitive arguments to evaluate among those possible advantages, and list the possible theories of victory to consider.

In planning for victory, it also becomes clear that weaker adversaries have their own theories of victory. American leaders should be cognizant of which theories of victory rivals such as China might pursue against the United States. Those include Cost Imposition as AI expenses continue to grow, Denial by producing good enough systems to minimize the U.S. benefits, and Devaluing by degrading markets on which U.S. companies or diplomats might otherwise hope to capitalize.

By aligning investment and effort behind AI's most substantial competitive advantages, by focusing on the most compelling theories of victory, and by mitigating against rival theories of victory, U.S. policymakers can increase the odds that AI will achieve its substantial promise.

Author

Andrew Lohn is a senior fellow leading the CyberAI Project at CSET.

Acknowledgments

The author would like to thank Igor Mikolic-Torreira, Neil Thompson, Cara LaPointe, Catherine Aiken, and Owen Daniels for helpful discussions during earlier drafts.



© 2026 by the Center for Security and Emerging Technology. This work is licensed under a Creative Commons Attribution-Non Commercial 4.0 International License.

To view a copy of this license, visit <https://creativecommons.org/licenses/by-nc/4.0/>.

Document Identifier: doi: 10.51593/20250024

Appendix A: Competitive Advantages

This appendix provides examples from familiar contests or game theory that illustrate how each of the eight types of competitive advantage can be decisive, substantial, negligible, or counterproductive. All possible future applications of AI cannot be listed here, so the hope is that these examples will provide simple analogies that can help decisionmakers reason through the various impacts of AI.

Information Asymmetry

An information asymmetry may exist if a competitor has access to more information, or if the competitor can make better sense of that information. AI is well suited to these information tasks but also faces limitations.

Decisive

Just as “any sufficiently advanced technology is indistinguishable from magic,” a sufficiently advanced AI could discover or invent things that a lesser adversary could not even imagine.⁵⁰ That level of intelligence overmatch is aspirational, especially given the advancing capabilities of peer competitors and open systems.⁵¹ However, smaller and more narrow information asymmetries can still lead to overmatch.

For example, if there are critical battles, supply chain weaknesses, or cyber vulnerabilities, then knowledge of those single-point, or few-point, failures can be decisive. Although decisively winning a battle does not guarantee victory in the war, winning many important battles increases the odds of victory. For AI to turn single-point failures into decisive victories, those single points need to exist, be identifiable, and be actionable. Adversaries are strongly incentivized to limit their number of single-point failures and to defend them. Given those defensive incentives, decisive information asymmetries may be rare.

Substantial

A famously important information asymmetry is World War II's Enigma machines and Alan Turing's early computers. In Operation Ultra, the British were able to decode German messages, using the premier artificial intelligence of the day to extract encryption keys. The resulting information asymmetry could be used decisively in battle, but because the British used it sparingly to avoid alerting the Germans and losing that advantage, the strategic impact of Enigma and Ultra on World War II remains a topic of debate.⁵²

Similarly in athletics, coaches and players use signs to communicate with one other in ways that can be observed and decoded by their opponents. The Houston Astros baseball team and the University of Michigan football team illicitly developed the ability to decode opponents' communications, which provided advantages.⁵³ While AI might provide advantages such as these, it is worth noting that simple low-tech solutions can nullify those advantages. Athletes can use radios or huddles, and the Germans could have changed their encryption schemes. In other words, AI-provided information advantages can be substantial, but they can also be fragile, short-lived, and non-decisive.

Negligible

If the information is not actionable, or if a competitor is overmatched for other reasons, then additional information is of little benefit. Continuing with the parallel to athletics, the basketball player Larry Bird is famous for telling his opponents what he would do before doing it. Similarly, military attackers often intentionally “telegraph” their attack to minimize collateral damage, knowing that defenders cannot prevent the strike despite the additional information.

In cybersecurity, Kerckhoffs' principle states that a system's security should not rely on keeping the approach secret.⁵⁴ Cyber defenders should be able to give attackers the exact blueprint of the system without worry. Of course, Kerckhoffs' principle is not always realized. If the system is poorly constructed, attackers can discover and exploit critical vulnerabilities, revealing information that does provide a substantial or decisive advantage.

Counterproductive

Additional information can also lead to negative outcomes. International-relations theory suggests that pairing superior knowledge with an incentive to misrepresent that information can lead rational nations into war.⁵⁵ That secret information can prevent competitors from reconciling their differences peacefully. For example, one party may know for certain that they will win in a fight, but only if they do not explain their plan to the opponent. The weaker competitor, believing it still has a chance, may carry on belligerently, while the stronger competitor has little incentive to appease it. A superior AI that provides these types of fragile confidences could lead to worse outcomes.

Perhaps the most famous game theoretic exploration of information asymmetries is Akerlof's Market for Lemons.⁵⁶ In it, used car salesmen can accurately estimate the

value of vehicles, but the less-skilled buyers cannot. Since buyers are not willing to pay top dollar without being able to verify a car's quality, they bid lower prices. With knowledgeable sellers unwilling to part with high-quality vehicles at low prices, the market degrades, leaving only dealers who sell low-quality vehicles ("lemons"). AI systems that are trustworthy to one group and not to another may create impasses that eliminate mutually agreeable compromises.

Conclusion

AI is certain to provide information, either directly or by helping to make sense of overwhelming floods of data. The resulting information asymmetries can be decisive if they reveal few-point failures that can be acted on. However, in many competitions, such as in telegraphed missile strikes or unguardable jump shots, additional information that AI may provide will do little to change the outcome. In others, such as intercepted communications, competitors will have the opportunity to adjust, making AI-provided advantages fragile and short-lived. Worse than fragility, information asymmetry can impede conflict resolution, whether in war or used car sales.

First Mover Advantage

AI promises superhuman speeds in the ability not only to act autonomously, but also to accelerate human decision-making and business processes. That initiative can provide a first mover advantage.

Decisive

First mover advantages can be decisive in races. They can also be decisive if the first competitor can lock in its preferred position, or if an early lead provides a compounding benefit.

High-frequency trading is a race, as are some high-speed cyber vulnerabilities known as race conditions.⁵⁷ These examples, and many like them, typically depend more on transmission speeds than on decision speeds. Where decision speed is indeed the limiting factor in winner-take-all races, then AI could prove decisive.

Another race is to develop patents. Current patent policy grants enduring exclusivity to the first inventor to file, although that does not necessarily lead to monopolies in practice.⁵⁸ A small lead can become a decisive advantage if a process or law such as patenting allows a competitor to lock in its position.

Other lock-in effects are not set by policy. Metcalfe's Law describes situations where being first gives a small advantage that compounds over time, becoming difficult to overcome.⁵⁹ Having more users in a social network attracts users to that network, making it easier to attract even more users. A similar effect has been proposed for AI systems where a marginal performance lead attracts more users, providing more data to further improve the system.⁶⁰ Or, alternatively, some propose that an AI system could recursively improve its own code, turning a small lead into a large lead.⁶¹ In these ways, Metcalfe's Law is a theoretical path to decisive first mover advantage, but MySpace, Friendster, Lycos, and AltaVista would all question it in practice.⁶²

Substantial

Aside from the Dictator and Ultimatum games where first movers demand whatever they wish, Stackelberg games are the most studied examples of first mover advantage. Two companies in the same market must decide how much to produce, and Stackelberg showed that the first to decide wins a larger (but not dominant) market share.⁶³ That advantage only holds, however, if the first mover is unable to adapt to the second mover's decisions. If it is inexpensive for the companies to adapt their production levels, then decisions about a starting point make little difference, and they split the market evenly. Theoretically, then, AI that improves adaptability, reducing the cost of change, can reduce or eliminate first mover advantages.

As another example, in chess, white moves first and wins about 52–56 percent of games.⁶⁴ Errors throughout play degrade that initial advantage, so the first mover advantage is small among weak players or in rushed games where errors are more common. At the highest levels, though, the first mover advantage is also small, and draws are common.⁶⁵ If perfect play in chess ends in a draw, as it does in checkers and tic-tac-toe, then AI advances would reduce or eliminate the first mover advantage. In competitions where strong play does not end in a draw, then AI may help competitors avoid blundering their initial advantage away.

Negligible

The business literature is filled with examples of companies that moved first and did not win the day.⁶⁶ And even in races, not all victors need to be first. A hiker does not need to outrun the bear, only the slowest hiker in the group. Similarly, a stock trader does not need to be the first to reap the rewards, only to move before the market does. In cases such as these, a faster AI system provides negligible additional value compared with a fast-enough system.

Counterproductive

A first mover advantage can also incentivize mutually detrimental behavior. That is clear in the Centipede game, where accepting the first mover advantage minimizes total winnings. Further, in a military context, viable first strike capabilities or an incentive to mobilize early can be detrimental to crisis stability. World War I is an infamous example, in which perceived first mover advantages rapidly turned an assassination into a global conflagration.⁶⁷ Incentivizing speed or initiative reduces time for negotiation and opportunities for human or bureaucratic oversight.

Initiative can also come at the expense of quality, safety, or security.⁶⁸ It can drive AI developers to behaviors such as excessive spending and resource consumption that neither they, nor broader society, would prefer. For example, the Stackelberg first mover advantage results in higher prices for consumers.⁶⁹ Therefore, policies that encourage racing by, for instance, promising to provide lock in effects should be carefully considered or reconsidered in light of how they affect risks and externalities from AI development and deployment.

Conclusion

AI may provide the speed to win races, help lock in advantages, or compound early benefits while helping to avoid blundering away an initial advantage. But the rush to move first can also exacerbate safety or reliability risks, and inflate the financial and environmental costs of AI development while reducing crisis stability, oversight, and de-escalation off-ramps. First mover incentives might also drive AI, or humans in the presence of AI, to logically choose mutually worse outcomes in collaborations and Centipede games. On the other hand, AI advances can reduce first mover advantages by improving responsiveness that unravels Stackelberg advantages, or, in chess, by helping black play to draw.⁷⁰

Second Mover Advantage

The early bird gets the worm, but the second mouse gets the cheese. If AI provides the ability to react quickly, then it can enable second mover advantages.

Decisive

“What can be seen can be hit” is not just Army lore, it is wisdom that applies in many contexts, including TV game shows.⁷¹ In a pop culture context, contestants on *The Price is Right*, who guess the price of consumer goods “without going over,” routinely

bid one dollar above their predecessors, all but eliminating the first mover's odds of success. As another example that draws more on AI's speed, robots can beat humans in Rock-Paper-Scissors by observing the human choice and selecting the winning option faster than humans can observe or react.⁷²

Moving first can also reveal weaknesses. Ahead of World War II, France invested heavily in fortifications across the Maginot line, but the Germans bypassed it easily due to a weakness near the Ardennes.⁷³ More recently, hundreds of articles describe cybersecurity's Maginot line. Digital systems are difficult to fortify in ways that do not leave weaknesses for attackers, as second movers, to discover and exploit. If AI can help identify these digital or fortification vulnerabilities, then it can empower second movers.

Substantial

Bidding second is not always as decisive as in *The Price is Right*. In negotiations with offers and counteroffers, the second bidder often does better in practice, but it is not often winner-take-all.⁷⁴

Moving second can also be advantageous in technology development. As training and development costs fall, the last company to make an AI system will have the least expenses to recover.⁷⁵ That is especially true if AI developers can draw from the first mover's products to facilitate their own development.⁷⁶

Negligible

Bidding once without going over is not the fairest way to run an auction, and there are many approaches that can reduce or remove the second mover advantage.⁷⁷ The more AI provides an advantage, the more likely disadvantaged competitors will be to establish or require fair processes.

AI might also reduce second mover advantages by creating better first moves. If *The Price is Right* contestants could use AI to place exactly correct bids more often, then moving second would become a weakness rather than an advantage. Similarly, if AI improves first bids in negotiations, then second bidders might lose their advantage. And if AI can do a better job of securing digital systems before they are released, then attackers will have less benefit from moving second.

Counterproductive

If all parties are incentivized to move second, then none will be compelled to take initiative, leading to deadlock.⁷⁸ In the canonical example of deadlock, four cars all arrive at an intersection at the same time. Each has to cede the right of way to another, so none are free to go. Deadlock is relatively rare in practice today, but it could become more common among AI agents that each wait rationally in their own best interests.

Conclusion

Second mover advantages may be less apparent than first mover advantages, but they are common nonetheless. AI may help second movers find flaws in physical or digital Maginot lines and outbid or out-negotiate competitors. AI may also provide the responsiveness to win in any situation where the best move becomes obvious once the first is played. In AI development itself, second movers can cut costs and need to innovate less. But AI could also weaken second mover advantages, for example by improving initial bids. And AI could increase the odds of deadlock by incentivizing all competitors to wait.

More Options

This section discusses the competitive advantage that could result if AI provides competitors with more options from which to choose. AI may help expand a playbook by accelerating the process of generating or validating options. It may also make employees or service members more broadly capable, giving them a wider range of viable taskings. And AI may be an additional option itself as an alternative to human or automated systems.

Decisive

One simple way to think about the addition of a disruptive highly capable new option is by analogy to children who upend the game Rock-Paper-Scissors by playing Dynamite.⁷⁹ Adding Dynamite is decisive only if it beats everything, and if only one player has it. If Dynamite beats Rock, Paper, and Scissors, and if only one player has it, then that player always wins. If both players have it, then the game always ends in a Dynamite-Dynamite draw. But as long as the new option has some vulnerability to exploit, it is not decisive.

Substantial

To see the effect of vulnerability, consider Dynamite that beats both Rock and Paper but can be cut by Scissors. If only one player has Dynamite, then that player only wins one extra game in every nine, boosting the win percentage from 50 percent to 56 percent.⁸⁰ In terms of AI, as long as AI systems are weak in some areas, disadvantaged competitors should be expected to adapt to exploit those weaknesses.

Chess also contextualizes the value of additional moves by scoring individual pieces. Chess is a drastic oversimplification of competition in business or war, but it is at least a familiar one.⁸¹ In it, Pawns represent one point each. Bishops and Knights are valued at three points for reaching an intermediate number of squares. Rooks reach more squares, for a value of five points, and Queens reach the most for nine points. If AI could enhance the capabilities of employees in the way that Pawns can be exchanged for Queens, then it could provide quite an advantage.

In the more applied cyber setting, attackers and defenders are constantly adding new moves as they develop new tactics, techniques, and procedures.⁸² These tactics are rarely decisive, especially if the adversary has an opportunity to adapt to any initial surprise. In the cat-and-mouse game of cybersecurity, and also more generally, adaptive adversaries are likely to iteratively gain and lose small advantages as a result of AI.

Negligible

Revisiting the Rock-Paper-Scissors example, as the game grows beyond just three moves, the value of new moves decreases.⁸³ That value can become so low that the new addition provides approximately zero value, even if it is much more capable than any other move. Like chess, Rock-Paper-Scissors is a dramatic oversimplification of real-world competitions, but it shows how weaker competitors can adapt to new highly capable AI-enabled actions in ways that limit their value.

The value of Dynamite can also be precisely zero if the additional move is not useful. Consider Dynamite that beats Rock but loses to both Scissors and Paper. In that case, Dynamite should never be used by either player, and it provides no value. In the current push for modernization and AI hype cycle, AI is often proposed or pursued for tasks that already have satisfactory solutions that are well-tested, reliable, and well-understood.⁸⁴ There are many cases where the addition of AI provides little benefit and introduces risks that should not be accepted.

Counterproductive

Even if an additional move or new technology adds little or no value, it can upend the competition or make prior approaches obsolete. The Rock-Paper-Scissors-Dynamite example also illustrates how prior approaches, and perhaps the technologies or humans that performed them, can become obsolete. That can happen even if the new technology is not wholly superior and does not lead to improved outcomes.⁸⁵

An increased number of moves can also lead to decision paralysis. And, similarly to decision paralysis, adding options can increase complexity. Increased complexity can potentially change the competition from one that is well-understood and solved to one that is too complex or chaotic for good decision-making.

Conclusion

AI may turn Pawns into Queens, increasing the versatility and capability of employees and soldiers. That may be enough advantage to secure victory, but competitors will adapt to limit those advantages and exploit the weaknesses in AI systems. Adapting to the addition of highly capable Dynamite in a game of Rock-Paper-Scissors can make shrink the advantage that Dynamite provides. In some cases, those adjustments can also have deleterious consequences such as obsolescence and unemployment, or can increase uncertainty and complexity, leading to worse decisions. In the push to modernize and cash in on new highly capable technology, the benefits may not always be as large as they seem, while the risks can be substantial.

Fewer Options

Perhaps counterintuitively, having fewer options to choose from can also be an advantage. As examples in an AI context, agents may benefit from not being given sufficient approval or authority to conduct a military retreat, to spend more than a capped budget on airfare, or to reimburse a complaining customer.

Decisive

Thomas Schelling famously described “the power to bind oneself.”⁸⁶ In Washington, DC, crosswalks become a game of Chicken when pedestrians, often wearing headphones, avert their gaze and march confidently into the crosswalk, knowing that by removing their own ability to respond, drivers have to wait. In higher-stakes competition, *Dr. Strangelove*’s fictional Doomsday machine automatically launches

nuclear weapons, and the nonfictional Soviet Dead Hand was built to automatically delegate authority in the event of a nuclear attack.⁸⁷

These decision rules, if communicated clearly to all competitors, can be decisive. An AI brinkman can achieve goals by being unable to coordinate on less favorable terms. That could include prescribed guidance for how to react in defeat, for instance whether to escalate, retreat, or self-destruct.

Substantial

More often, the limitations on AI agents will be less constraining. The limitations will probably be initially or primarily intended as safeguards to protect against concerns about AI's reliability or security. AI agents cannot be fully trusted, especially given how easily they can be manipulated in an adversarial context.⁸⁸ But those same safeguards can also be crafted for competitive advantage.

Consider bargaining with limits. An AI agent with a user-specified price cap has additional leverage when negotiating a block of hotel rooms. In negotiations among AI agents, it may be the weaker of the AI agents that requires the most stringent safeguards and that gains the most negotiating leverage.

Negligible

As in the previous section, which showed cases where adding options provides little or no value, removing those same unimportant options also provides little or no value.

Counterproductive

Although perhaps decisive, the earlier example from *Dr. Strangelove* can also be counterproductive, as made clear in the movie. The same could be true in lower-stakes competitions where decisions are pre-scripted toward escalation or self-destruction, potentially leading to mutually undesirable outcomes that might have been avoidable with more options.

As AI agents and systems are built with constraining guidance or hardcoded rules, developers and deployers should be careful to consider the potentially complex multi-agent dynamics as well as likely adversary actions and countermeasures. That is especially important in highly interactive and autonomous cases such as those that occur in the military, cyber, and possibly financial domains.

There are also intuitive disadvantages to having fewer options. Having fewer options implies less flexibility and adaptability. It also implies more predictability. In most cases fewer options are probably a disadvantage, and in most cases, AI will probably provide more options rather than fewer.

Conclusion

Whether booking hotel rooms or fighting to the robot death, a more limited AI system or delegate has some negotiating leverage that can provide an advantage. In stark cases, such as *Dr. Strangelove's* Doomsday Device, that advantage could be decisive, but the risks of brinkmanship are typically not worth the benefits, and AI typically increases adaptability more than restricts it. The limitations placed on AI agents will probably be intended primarily as safeguards against accidental errors or malicious manipulation, but users will certainly adjust those safeguards to gain additional advantage. And in this case, the additional advantage from more restrictive safeguards might go to the weaker AI systems that need them more, or that can justify them without appearing to be duplicitous.

More Time per Turn

Many of the effects of time have already been covered in earlier sections. Having more time, or having the speed to make better use of time, can help win races, as detailed in the section "First Mover Advantages." Time can also provide more responsiveness or adaptability, discussed in the section "Second Mover Advantages." And more time can provide more options if it allows for more actions per minute, or if it accelerates the process of developing options such as identifying alternatives, soliciting feedback, or anticipating outcomes.

Without reiterating the benefits of time that can be cast to the advantages previously discussed, here the focus is on more ways in which time can be helpful.

Decisive

A simple and common framework for considering the stages of competition is the OODA loop.⁸⁹ It was developed in the U.S. Air Force, where time pressures are significant. Fighter pilots often talk about how an accelerated ability to move through the OODA loop can be decisive. In a more general sense, moving rapidly through the OODA loop is probably more useful than decisive.

Substantial

Competitors can use their time to outwait opponents. For example, politicians can filibuster, footballers can feign injury toward the end of games, and laborers can strike. Consuming the available time can be a risky choice that AI may be able to help facilitate. An AI-enabled competitor may be able to delay until the last minute if AI can reduce the uncertainty in how long a process will take, or if it shrinks the time from initiation to completion.

Tactics to manipulate time and delay are already prevalent. For example, defenders set rate limits for downloads or impose manual or multifactor confirmations before granting access. The stock market has circuit breakers that pause transactions during suspicious or dangerous market behaviors. And in the opposite sense, time pressure is increased by setting expiration dates such as for stock options, offers in negotiations, computer logins, cryptographic certificates, or clocks in chess. As AI is deployed more frequently, simple methods such as these to impose or remove time advantages will probably become more common. The cyber and stock market circuit breaker examples are already the result of automation.

Negligible

Time matters only if there is not enough of it. If a competitor has sufficient time to make the best decision, then additional time provides no advantage, and having lost time provides no disadvantage. AI can make time factors irrelevant when it can enable competitors to arrive at their best possible decision more quickly.

Counterproductive

Incentivizing faster models or shorter timelines risks introducing mistakes. And consuming time for competitive advantage can force adversaries to act quickly. Leaving little time for adversaries to respond forces them to make rash decisions, which could be worse for all sides.

Conclusion

AI might accelerate the stages of the OODA loop for more than just fighter pilots. And it might help eat up the clock for more than just footballers with a lead to protect. But if AI helps make decisions quickly, then adversaries may find their remaining time perfectly sufficient. Rather than consume time, AI may be an impetus for instituting new time controls, as it has been for throttling download speeds and stock market

circuit breakers. Carefully constructed, those time controls can help incentivize good decisions rather than forcing rash ones.

More Material/Materiel

A longtime maxim in military circles is that “quantity has a quality of its own.” More employees can achieve superior results with more resources, whether pens or swords. Quantity can also beget quantity. If AI enables more production or lowers costs, the proceeds can be reinvested to further increase production.

The fictional Colonel Blotto deals with quantitative trade-offs in theoretical games about military scenarios, election campaigning, and many other domains.⁹⁰ In the game, Colonel Blotto has a total number of armies, drones, or dollars to allocate across various battlefields without knowing how much the adversary will allocate to those battles. A player wins battles in which the allocation exceeds the enemy’s, and loses where it is lower.

Decisive

A competitor can be overmatched by materiel. In a Colonel Blotto game, a competitor could outnumber opponents in enough fields of battle to guarantee victory. The necessary ratio for overmatch may not be one-to-one, as in the basic Colonel Blotto game; a common rule of thumb in military contexts is that three attackers are needed per defender.⁹¹ That three-to-one rule is certainly not definitive, but it has proven surprisingly reliable for centuries and through technology transitions, including in recent Russian and Ukrainian use of fighting vehicles.⁹² In AI-enabled roboticized battles of attrition, Colonel Blotto’s quantitative approach may be more appropriate than ever.

Substantial

In an athletics context, small quantitative advantages can have substantial effects. Some studies show that reducing a soccer team to 10 players early in a game nearly doubles the opponent’s win probability, from 37.5 to 65 percent.⁹³ Other research suggests the advantage may be much smaller or even nonexistent.⁹⁴

In a cyber context, differences in scale can be more dramatic to lesser effect. Botnets collect large numbers of devices that can be used to overwhelm defenses in Distributed Denial of Service attacks. DDoS attacks are common, but they are also often defended successfully.⁹⁵ The defensive task is more difficult against more

sophisticated attacks, but defense-in-depth strategies can hold up against very large numbers of attackers, even if AI makes DDoS attacks more tailored and creative.⁹⁶

More humorously, the whole of society was recently captivated in debate about whether 100 human men could defeat a single gorilla in combat.⁹⁷ The 100:1 ratio is an indication of at least the perceived value of a material advantage. That same ratio comes up in a macroeconomic context. A country with few laborers who are productive can outproduce a larger country that is inefficient. The most productive countries produce \$100 of output per laborer-hour, whereas output per laborer-hour is approximately \$1 for the least efficient countries.⁹⁸

Economists have many simple and complex models of productivity and scale.⁹⁹ In such models, AI can be expected to change many different parameters. It can increase the number of tasks that can be performed, the productivity of workers, or even the number of workers if AI replaces workers.¹⁰⁰ It is still unclear whether AI will narrow productivity gaps by aiding low-skilled labor or exacerbate those gaps by empowering high-skilled labor. It is also unclear whether AI will supplement small labor pools, providing them with large digital workforces. And more broadly, it is unclear how much AI will affect production overall. Solow's paradox may apply to AI: AI may revolutionize many aspects of life and work without having much effect on productivity.¹⁰¹

Negligible

There are many possible explanations for Solow's paradox in AI that suggest AI might not dramatically increase production. Some hypotheses are that the importance of technology might be overrated, that all competitors could advance together without improving anyone's relative capability, that it takes longer to realize gains than anticipated, or that once a task becomes automatable, it declines in value.¹⁰²

Apart from whether AI will provide quantity is the question of when quantity provides advantage. In parallel computation, having more computers increases the amount of work that can be done. That can be extremely effective and arguably launched the current AI boom. But in many applications, there are bottlenecks where many computers have to wait for the results of one, eliminating the advantage of having those computers. This concept, known in computer science as Amdahl's Law, is familiar in many fields.¹⁰³ Bottlenecks are common in nearly all domains and can limit the advantages of increased scale that AI promises to provide.

Counterproductive

Perceived advantages of scale can lead to arms races. The Cold War idea that each country needed more nuclear weapons than its adversary in order to maintain a reliable retaliatory capability led to massive build-outs. Those build-outs were threatening, expensive, and created other proliferation and accident risks. Similar risks exist today for AI-enabled weaponry such as drones that fight by attrition, or among AI developers who are trying to build models or systems that outperform those of their competitors.¹⁰⁴

Scale is also costly. AI is expensive in terms of dollars spent and resources consumed, both to develop and to operate.¹⁰⁵ If AI ultimately provides little realizable value, profit, or competitive advantage, then its resource consumption could be regrettable.¹⁰⁶

Conclusion

Colonel Blotto might claim that quantity's quality of its own can be decisive. AI might enable massive drone swarms or help to produce enough goods or capital to overwhelm an adversary. But superior numbers do not necessarily provide superior might. It may be that a military's defensive 3-to-1 advantage, a gorilla's 100-to-1 advantage, or a cyber defenders' layered barriers set too high a bar for AI-enabled material advantage to be decisive.

Further, both Gene Amdahl's and Robert Solow's observations might continue to hold true if AI development or use cases run into bottlenecks or fail to produce the expansions that many anticipate. If so, the costs of AI development may be regrettable. But either way, the racing that has already begun may have unfortunate consequences.

Differing Objectives

The advantages discussed so far implicitly assume that competitors are trying to achieve similar goals. Another type of advantage comes from having an easier objective than competitors. It may be easier to repel than to invade, or it may be easier to hide than to seek.¹⁰⁷ In such games, the environment can matter more than the skill of players, and AI is likely to change both in many cases.

Decisive

AI has already changed the balance in many hide-and-seek games, laying bare some haystacks that used to hide needles. Facial recognition has advanced so that people may now stand out in a crowd where they used to blend in, providing decisive victories

to some seekers. In other cases, the haystack is as impenetrable as ever. An automated system can guess many passwords, but a well-chosen password should be effectively impossible to guess, even for arbitrarily advanced AI.

Substantial

AI advances typically help the searcher, but hiders can often adapt to weaken the searcher's advantage. People who want to remain hidden can evade facial recognition with masks, or they can approach a military checkpoint in cardboard boxes.¹⁰⁸

And where a hider may feel secure, AI may alter the game to reduce that security. In the password example, an attacker does not need to guess the password if there is a flaw in the encryption software. If AI can help break either the encryption algorithm or its implementation, then it can reveal the hidden information or discover an encryption key in fewer guesses.

Negligible

As AI makes the game easier for either hiders or seekers, each adapts to compensate. Cyber attackers and defenders are in a constant cat-and-mouse game, holding the overall balance more or less steady. Defenders have added multi-factor authentication, encrypted websites, and built layers of antivirus and anomaly detection, but breaches are still common. And attackers have increased their use of zero-day vulnerabilities, ostensibly as defenders have gotten wiser and harder to defeat by more common techniques.¹⁰⁹ That cat-and-mouse balance also applies to military intruders and repellers in the literal trenches of modern warfare.¹¹⁰ As AI provides advantages that make one objective easier, those pursuing counter-objectives will often have opportunities to make simple changes to compensate.

Counterproductive

As AI alters these search or pursuit-evasion games, the impacts extend to regular citizens who are not part of the competition. While better search helps to bring wrongdoers to justice, it also introduces privacy and security concerns for law-abiding citizens who need the freedom to dissent and the anonymity to act freely. And while better evasion helps to protect privacy, it also enables illicit actors to demand ransoms for cyberattacks or to distribute child sexual abuse material.

Advances in search can also be counterproductive for the searcher. For example, if a nation can locate an adversary's nuclear weapons, then the vulnerable nation might move to a launch-on-warning posture.¹¹¹ That would put the more advanced nation in

a more threatened position. Alternatively, the vulnerable nation might turn to deceptive tactics to which AI systems are especially vulnerable.¹¹² That could lead to more miscommunications and miscalculations, which increase the risks and uncertainty for everyone involved.

Conclusion

Advances in AI can help find the needle in some haystacks, but it is often easy enough to increase the size of the haystack or adapt in other ways to compensate. Throughout technology transitions, competitors in cyber and trench warfare have made those adaptations to maintain a balance. Such shifts can sometimes affect others outside of the main competition, as both dissenters and criminals have learned. The risks can also rise for those who would seem to gain an advantage. For example, everyone is worse off from incentivizing nuclear launch on warning, and from incentivizing deception that can lead to mistakes and miscalculations.

Appendix B: Walkthrough of AI Superiority to Victory in Construction Materials Industry

This appendix presents an example of how AI might be expected to have an impact among competitors in a construction materials industry such as steel or concrete manufacturing. It describes how an AI superiority could develop, the competitive advantages that superiority could provide, some likely countermeasures to anticipate, and which theories of victory would be reasonable to pursue or to anticipate from competitors with inferior AI.

This is not a complete or thoroughly conducted analysis. It is meant only to illustrate the types of arguments that can be made throughout the chain of reasoning for which this report advocates.

Construction Materials Industry such as Steel or Concrete

The industry of construction materials production using steel or concrete is well developed, with large firms that have established clients, market shares, business processes, production processes, and capital infrastructure. The assumption here is that firms in this industry tend to have relatively few large contracts that last for years and are negotiated over a period of months or more.

AI Superiority

Companies within major manufacturing are unlikely to develop their own frontier AI models, as the expense is far too large and the expertise too distinct from their main business. A firm could invest in digitizing production or business processes to better incorporate AI tools, or to develop tailored datasets. It could then further invest in a small team of AI experts to post-train an open model on its data to create a proprietary model that is specialized for its purposes. Alternatively, it could team with a frontier model provider that could post-train a frontier model using that firm's data. In that case, it would have to trust the model provider to safeguard its secrets.

These tailored models could be related to the firm's production processes to improve research and development and develop new materials and processes, or they could help to improve the efficiency and reduce cost for existing products. The models could also be applied to internal business processes such as logistics or personnel management, reducing costs or improving efficiency in indirect ways. And models could assess the broader industry to better understand the positions of competitors and clients, allowing for more favorable negotiations and contracts.

A firm's AI superiority does not necessarily require access to superior AI. It could invest in more pilot projects earlier, or adapt its processes to incorporate AI to a greater degree. This still requires initial investment in pilot projects, digitization, and process changes.

Competitive Advantages

This section briefly presents examples of possible competitive advantages that a firm could expect from the AI superiority described above.

Information Asymmetry

- An AI leader might develop an edge in research and development, allowing it to produce more advanced materials, or to develop more efficient production processes.
- An AI leader might develop a better understanding of the market for both customers and competitors, allowing for more favorable negotiations and contracts.

First Mover

- An AI leader might be able to write patents for its materials or processes to lock in early leads in research and development.
- New products could build on old products, thereby gaining an edge that provides a step toward future improvements, with followers always a step behind. This is a bit speculative and requires several uncertain conditions to hold.
- An AI leader might be more responsive to customers, offering bids faster than competitors.
 - This might not be a major factor in this sector, where negotiations are expected to occur over months. Customers are likely to wait for slower competitors to provide competing bids.

Second Mover

- This is an established industry, so many of the first and second moves, such as production capacity and pricing, have already been made.

- An AI leader may be able to make more well-informed counteroffers or bids by incorporating more information about the market or better understanding its own business processes and adaptability. For example, it could assess a wider range of scenarios.

More Options

- If AI uplifts employees, an AI leader may be able to offer a broader range of products than its more limited expertise would allow without AI assistance.
 - This may not be true if the firm's offerings are limited by other factors such as processing equipment or other high-capital infrastructure, or if the business focus is intentionally aimed at a narrow market.

Fewer Options

- This is probably not a large factor for this industry. Most critical decisions are important enough for humans to be involved. Humans would use AI as an aid, rather than fully delegate authority to AI that can be capped or limited.
 - This would be different if the industry were driven by a large number of individually small negotiations, where automated decisions were needed and constraints on those decisions could be critical.

More Time

- This is probably not a large factor for this industry because negotiations are expected to proceed on human timescales, and customers are willing to wait for competitors to bid. There are not many time-limited move-countermove dynamics.

More Materiel

- AI in the form of roboticization could allow production to run around the clock.
 - This presumes that processes are not already running at their maximum. Processes might already be limited by bottlenecks in chemistry or physics such as cooling or mixing rates, or human employees could already work around the clock in shifts.

- An AI leader might be able to free up more resources if it could reduce employee costs, by either having fewer employees or hiring lower-expertise employees at lower salaries.

Differing Objectives

- This is probably not a large factor in this industry, where firms are expected to have the same profit-maximizing objectives.
- At a more nuanced level, firms could compete to provide different types of value. For example, a firm could prioritize customer responsiveness instead of product quality or cost.

Competitor Countermeasures

Competitors with inferior AI access or deployment can adapt or take countermeasures to limit their disadvantages, some of which are described here.

- All companies have access to AI. The gap in performance between specialized AI trained on data tailored to the industry and AI trained to a particular firm may be small by comparison to out-of-the box frontier models or open models.
 - Proprietary data developed by one firm may be difficult to keep from leaking into training pipelines or being stolen or distilled by rivals.
- Companies may also reap the benefits of another's AI investment by reverse engineering the leader's products or poaching its talent.
 - Reverse engineering could be especially problematic if competition included firms outside the jurisdiction or enforceability of patent laws, such as in rival countries.
- A company that feels threatened by another's AI-enabled market analysis can seed the information environment with deceptive data that misrepresents the market, its relative position, and its strategies.

Theories of Victory

Considering the possible competitive advantages and competitor countermeasures, some theories of victory seem more promising than others.

Cooperation

- Companies across the industry could pool their data and learn from one another's pilot projects to create better products produced more efficiently for all.
- If AI adoption increases market size, then all companies could come out ahead even as they compete for market share.

Dominance

- Dominance based on model access alone does not seem likely because all firms will have access to comparable AI.
- Dominance could come from digitizing processes and incorporating AI more aggressively than others, but it is not obvious that the advantages would be decisive.
 - Production processes are likely to be limited by factors such as chemistry and physics timelines, access to raw materials, or the throughput of capital equipment.
 - Negotiations and contracting is slow and deliberate, nullifying AI's speed advantage.
- Dominance may be viable if roboticization, workforce cuts, and streamlined logistics/business processes create a sizable cost savings that can undercut competitors. That requires those aspects of the business to be inefficient enough without AI that improvements with AI would be dramatic enough to be decisive.

Denial

- It is not obvious how an AI leader would pursue denial against a rival.
- Low-AI firms could try to deny an AI advantage by muddying the information environment with disinformation about their bids, costs, ambitions, etc.
- A firm could try to negotiate with closed model providers for exclusive rights to its services in order to deny competitors access to frontier AI capabilities.

Devaluing

- Cost-cutting measures such as those in pursuit of a dominance theory of victory could squeeze margins out of the industry.

Brinkmanship

- It is not obvious how brinkmanship would be used in the context of AI, beyond normal competitive dynamics such as pricing and investment.

Cost Imposition

- A low-AI firm might try to reap the benefits of its AI-enabled competitor's digitization without making the AI investments itself. A low-AI firm may reverse engineer progress or poach talent, especially if patents are not enforceable.

Conclusion

An AI-forward construction materials firm might digitize and invest in AI pilot projects and process reform in pursuit of a Dominance theory of victory. It might succeed if it avoids bottlenecks that seem likely to limit its improvements, or if it sufficiently reduces operating costs. However, in the fight over cost reduction, an AI-forward construction materials firm needs to be careful not to devalue its industry too much, or fall into a trap of overinvesting in technology and advances that competitors can degrade with disinformation in a Denial theory of victory or adopt more cheaply in a Cost Imposition theory of victory.

Appendix C: Walkthrough of AI Superiority to Victory in Conflict Between Large and Medium-Sized Militaries

This appendix presents an example of how AI might be expected to have an impact among competitors in a conflict between a large military and a medium-sized one. It describes how an AI superiority could develop, the competitive advantages that superiority could provide, some likely countermeasures to anticipate, and which theories of victory would be reasonable to pursue or to anticipate from competitors with inferior AI.

The following is not a thorough or rigorously conducted analysis. It should not be used as the basis for decisions about military strategy or investment. Instead, the intent is to illustrate the process advocated in this report, and to offer the types of arguments that could be made.

A Large AI-Enabled Military Facing a Medium Low-AI Military

One competitor in this example is a large expeditionary military with substantial capability across domains including land, sea, air, space, cyber, plus a nuclear deterrent. This large military is also supported by substantial non-military or military-adjacent economic and industrial power at home. That economic and industrial power includes a thriving AI industry that provides capable AI models as well as the computing power and expertise to operate them at scale. The description of a large AI-capable military with a thriving AI industry naturally evokes images of the United States or China, but it could also include countries such as the United Kingdom, France, and Israel, which lag in model development but have the technical talent to integrate and utilize AI effectively. The larger military in this example could also be an alliance of militaries.

The smaller military has less capability across domains, with fewer ships, artillery, planes, and infantry. It has little to no space infrastructure, modest cyber capabilities, and no nuclear deterrent. Its ambitions are also smaller, engaging in conflicts only within its geographical region. This smaller military relies on open models or foreign AI providers and has more limited indigenous computing capabilities and technical talent for integrating AI at scale. Nonetheless, it sees AI as a way to make up for its other military shortcomings at low cost, such as through attritable drones, open-source intelligence analysis, and enhanced cyber capabilities. These could be countries such as Ukraine, Iran, or Taiwan.

AI Superiority

The AI superiority in this example comes in several forms, all to limited degrees. One form of advantage is having access to the superior models that the larger nation can create. That superiority is limited by the presumption that open models remain near the frontier and can be adopted by the smaller nation. Another source of superiority is the larger nation's computing power that allows it to scale AI more aggressively than the smaller nation. That is not much of a factor for embedded AI, such as in drones. It is also only a small factor in applications that do not require AI usage to scale, such as in generating a few AI disinformation videos or cyber exploits. When AI needs to scale, for instance in targeted disinformation campaigns or widespread open-source intelligence analysis, computing power can be a limiting factor. And finally, AI superiority comes from technical talent that can design processes and applications to take advantage of AI most effectively. This is less about AI development and more about its integration and usability.

Competitive Advantages

This section briefly presents examples of possible competitive advantages that a nation or military could expect from the AI superiority just described.

Information Asymmetry

- Most of the information asymmetry between these countries does not come from AI directly, but from superior data collection including remote sensing from space, cyber capabilities, and intelligence services that might include human intelligence.
- Overall, AI probably reduces the information asymmetry.
 - A highly capable intelligence agency is overwhelmed by data flows but has already learned to prioritize. AI allows it to include lower-priority data and assessments.
 - A fledgling nation can use AI for open-source intelligence or possibly to develop cyber espionage capability in ways that are a categorical improvement on its ability without AI.
 - A possible counterargument is that a fledgling nation does not collect enough data to make as much use of AI compared with a nation that has established intelligence collection. An actual analysis rather than this

example discussion should debate and resolve these types of counterarguments more fully.

First Mover

- AI enables many targets to be struck quickly, especially in the early stages of a conflict, before response actions and potentially slow materiel flows can limit operational speed. Increased scale of early strikes can have a disarming effect. This high rate of attack was seen in Israeli strikes in Gaza and U.S. strikes in Iran.¹¹³
- Delays are often caused by planning. For example, AI has been able to accelerate the bureaucratic aspects of sorties, getting planes in the air faster.¹¹⁴
- AI can allow attackers to accelerate targeting processes, from detection to damage assessment.¹¹⁵ Very short processes have long been possible with loitering munitions, high-speed anti-radiation missiles, or landmines, and are now possible on attack drones with cameras.¹¹⁶ On the other hand, fully automated kill decisions may provide only minimal acceleration beyond humans with a video feed.

Second Mover

- AI can enable rapid retargeting or adaptive targeting so that after a moving or hidden target is reacquired, it can be acted on faster.
 - This may be limited by physical constraints such as missile flight times.

More Options

- AI does not create many novel weapons; instead, it promises to make better use of existing ones. Guns, bombs, and artillery will likely remain the primary weapons of war.
- AI may make military platforms more versatile. A small drone may be able to collect intelligence, reprioritize collection, conduct electronic warfare, and perform kinetic strikes in ways that previously required larger, more expensive, and less plentiful platforms or remote human pilots.

- Human soldiers may also become more versatile. With less training, AI may help a single soldier manage medical fieldwork, electronic warfare, fires planning and operation, cyber operations, and more.
- This may provide more uplift to the smaller nation that was more limited before AI. Whereas without AI it might have had few cyber operators or trained pilots, the larger nation could already employ all these capabilities without AI.
- Much of winning and losing can involve options beyond those of traditional combat. It involves rallying international and domestic support. That can include economic actions such as sanctions against Russia, Taiwan's Silicon Shield, or closing the Strait of Hormuz.
 - AI might support or enhance these additional options for "politics by other means" through media generation, better monitoring of the information environment, or disinformation.¹¹⁷
 - AI may help provide more options by analyzing or gaming out complex economic or political decisions such as closing the Strait of Hormuz or blocking shipments of semiconductors from Taiwan.

Fewer Options

- This competitive advantage is not obviously substantial for this scenario.
- AI in automated battles or intelligence analysis may be limited or safeguarded, but the advantages those limits provide probably have a second-order effect rather than being a primary driver of advantage.

More Time

- First and second mover advantages have been discussed previously and are not included here, except to note that time may be unaffected or compressed more than expanded.
 - Many durations will continue to be bottlenecked by slow processes such as troop movements and missile flight times.
 - It has been proposed that increased decision speed could result in less time to make decisions because of a Jevons paradox for decisions. It may

be that faster decisions lead to more decisions or compressed decision times, not more time per decision.¹¹⁸

More Materiel

- Overwhelming numbers can be decisive in battles. As General John Pershing said, “Infantry wins battles, logistics wins wars.”¹¹⁹
 - Because large militaries have focused on exquisite and expensive platforms, it is not obvious that a large nation’s superior industrial capacity translates into a numerical advantage over smaller nations that produce cheaper platforms.
 - Large militaries are learning lessons from recent conflicts and creating low-cost options and concepts of operations.
 - AI may be useful for improving normal business processes or logistics. For example, logistics has been envisioned to improve predictions for ammunition resupplies, resulting in more materiel that can be used.¹²⁰

Differing Objectives

- The smaller nation is probably defending within its borders or region, perhaps giving it the historical three-to-one advantage if that rule of thumb persists through the AI era.
- Defense also has an automation incentive. Defenders are more justified and incentivized to automate more aggressively.¹²¹
 - There will be defensive impetus on both sides, but especially within the smaller nation that is acting within its borders or region.

Competitor Countermeasures

Competitors with inferior AI access or deployment can adapt or take countermeasures to limit their disadvantages, some of which are described here.

- All parties are incentivized to use deception and decoys in ways that decrease the various competitive advantages to expect from AI.

- AI is likely to continue to be more susceptible to deception than humans, if only because deceivers can test and revise their deceptions against machines billions of times.¹²²
- Deception can involve fake targets, traditional decoys, or patches tailored to deceive AI systems.¹²³
- Deception can also involve open-source disinformation, ruses, or data poisoning to affect more text-based intelligence collection and analysis or decision-support systems.¹²⁴
- Similar manipulations to the aspects just listed can intentionally cause AI errors, collateral damage, or fratricide in order to reduce support for AI integration or for the conflict more generally.
 - The public already, in general, has an unfavorable view of AI, especially in the United States.¹²⁵ Errors that are presumed to be due to AI, such as the Minab school bombing, could reduce support for AI integration and application, weakening the anticipated superiority or competitive advantages.¹²⁶
 - Domestic and international support for automation may be especially fragile for the nation that is seen to be on offense rather than the defense.¹²⁷
- Recognizing the first mover advantage from increased scale of initial attack such as was seen in Iran and Gaza, countries may harden or hide their assets more aggressively during peacetime or early stages of crisis.
 - Nations that feel vulnerable to first mover disarming strikes or second mover rapid retargeting may shift to a use-it-or-lose-it mentality. They may fire more aggressively or preemptively if they suspect they are at risk.
 - Nations may also adapt by relocating faster after fires in order to minimize the second mover advantage from AI accelerating rapid retargeting.
 - Nations may also partially randomize attack lines and timings to counteract AI's predictive logistics, including for where and when ammunition will be needed.

- Nations may try to create bottlenecks, for instance by forcing the fight to diverse geographic locations to which the other would have to relocate or else attack using expensive munitions.
 - Nations could plan for several different options for progressing through the conflict and select from them randomly, or select based on changing on-the-ground circumstances that they expect would be difficult for the adversary's AI to predict. Serializing decisions in this way would nullify some possible planning, logistics, and intelligence advantages.

What Does It Mean to Win?

Victory in this example is considered to be an improved positioning for a negotiated peace in terms of strategic goals related to territory, regime change, military restrictions such as nuclear abolishment, or access to trade, natural resources, or thoroughfares. A victor in this military conflict shifts the negotiations in its favor.¹²⁸ It is possible for both parties to lose more in the fight than they gain in the negotiation. A primary goal throughout conflict is to avoid escalation that imposes more costs than the eventual negotiated peace can justify.

Who Has the AI Advantage?

This example is constructed for AI superiority to belong to the larger nation with AI industry, talent, and computing power, but it is not obvious that the larger nation does have AI superiority. That is because the smaller nation has only a regional focus and is likely to be on the defensive—for which AI integration and use is more palatable. The larger nation may also have dwindling domestic and international support for AI in general, especially militarization of AI.

Further, in terms of turning AI into competitive advantage, AI may play an equalizing role even if access to it is not equal. The larger nation already has well-functioning intelligence agencies, cyber capabilities, electronic warfare, aerial surveillance, and attack platforms, along with well-trained infantry, medics, and logisticians. Top-tier AI may only slightly uplift such existing capabilities, whereas even limited access to second-tier AI may provide substantial uplift to a nation that does not already excel in those areas.

Theories of Victory

Considering the possible competitive advantages and competitor countermeasures, some theories of victory seem more promising than others. In planning for and

conducting conflict, a theory of victory should be based on the broader conflict, not focused primarily on AI. For the sake of this report, the focus is on how AI might factor into theories of victory.

Cooperation

- Cooperation is difficult to envision in this military conflict scenario. Still, the goal is negotiated peace, so cooperation is not only possible but required to some degree.
- It is even less obvious how AI directly supports cooperation, although it could support aspects of it such as translation, brainstorming mutually agreeable resolutions, and clarifying cultural sensitivities. These types of functions already have human experts, so the uplift from AI might be marginal.

Dominance

- The larger nation has a few possible paths to dominance in this scenario.
 - Overwhelming numbers from manufacturing such as drones or roboticized warfare
 - Overwhelming numbers from logistics such as predictive ammunition or personnel needs, and superior sortie counts
 - First mover advantages that enable disarming strikes
 - Second mover advantages with rapid retargeting that forces the adversary to stay hidden
 - This requires rapid retargeting to be faster than hiding, both of which can become limited by physical movement times rather than decision speed.
- The overwhelming numbers and first mover advantages require a level of aggression that has been unacceptable. Perhaps if AI increases precision sufficiently, aggressive preemptive counterforce strikes that limit collateral damage would become more palatable.

Denial

- The larger nation has a few possible paths to denial.

- An AI denial strategy for the larger nation can focus on limiting compute shipments to the smaller nation if it can control shipments. Compute capacity may already be limited in the small country, so further constraints may have little additional effect.
- The larger nation can strike datacenters to reduce compute access. This might have limited effect because much of the smaller nation's AI will be distributed in small clusters for inference or fine-tuning only, or embedded in individual weapons.
- The smaller nation also has a few possible paths to denial.
 - Degrade domestic and international support for the larger nation's militarization of AI.
 - Achieve greater uplift with inferior AI than the larger nation receives from superior AI. The larger nation may have superior AI but decreased competitive advantage from AI.
 - Integrate and authorize AI more aggressively as a result of being on the defensive.

Devaluing

- Both nations can attempt to devalue AI for military applications by:
 - using decoys, deceptions, ruses, and randomized behaviors.
 - seeding the information environment with false information or poisoning datasets and models so that predictions such as logistics or decisions about intentions or troop locations are misguided.

Brinkmanship

- Either nation, but especially the one on the defensive, can move to a more aggressive use-it-or-lose-it mentality.

Cost Imposition

- Cost imposition has recently focused on forcing adversaries to use high-cost munitions to shoot down low-cost attritable platforms such as Shahed drones.¹²⁹

- Drone or roboticized buildouts could lead to expensive arms races for battles of mechanized attrition. This may have little effect because much of the race is to build the least expensive platforms.

Conclusion

Each theory of victory, most of the competitive advantages, and even arguments for which nation has AI superiority are mixed, with points favoring and cutting against each side. This implies that AI will play a role in the conflict, but it will not be definitive or guide either side to victory. That said, the bar for victory is not necessarily high. Victory in this scenario is simply gaining more favorable terms for negotiated peace, not complete submission of an adversary.

In terms of theories of victory, it would be difficult for either side to deny AI capability, and both should be expected to devalue it with decoys and misdirection, but only to limited degrees. The most promising prospective theory of victory for the larger nation is dominance, where its manufacturing capacity may be able to provide overwhelming numbers or disarming first strikes. It is not a promising theory of victory from AI because overwhelming force was already an option for the larger nation without AI. The advent of AI changes little, unless precise mass is able to reduce collateral damage to such an extent that overwhelming force becomes palatable, and if defenders do not sufficiently manage to hide and deceive.

The most promising theory of victory for the smaller nation is cost imposition by forcing the larger nation to expend large amounts of expensive munitions in exchange for low-cost attritable platforms. This is also becoming less viable as a theory of victory as large nations increase their low-cost arsenals and improve counter-drone tactics.

By the low-bar definition of victory being merely an improved negotiating position, and in the absence of a compelling theory of victory, AI probably favors the smaller nation. That is because of the leveling effect in which AI gives the smaller nation low-cost aerial, naval, cyber, electronic warfare, and open-source intelligence platforms that they might not otherwise have. The smaller nation, which is likely to be on defense, also has extra incentive and support for militarized autonomy.

Endnotes

¹ The White House, *Winning the Race: America's AI Action Plan* (Washington, D.C.: The White House, July 2025), <https://www.whitehouse.gov/wp-content/uploads/2025/07/Americas-AI-Action-Plan.pdf>.

² Aditya Challapally, Chris Pease, Ramesh Raskar, and Pradyumna Chari, "The GenAI Divide: State of AI in Business 2025" (MIT NANDA, July 2025), https://mlq.ai/media/quarterly_decks/v0.1_State_of_AI_in_Business_2025_Report.pdf; Kristina McElheran, Mu-Jeung Yang, Zachary Kroff, and Erik Brynjolfsson, "The Rise of Industrial AI in America: Microfoundations of the Productivity J-Curve(s)" (CES Working Paper No. 25-27, U.S. Census Bureau, April 2025), <https://www2.census.gov/library/working-papers/2025/adrm/ces/CES-WP-25-27.pdf>.

³ Andrew Lohn, "Scaling AI: Cost and Performance of AI at the Leading Edge" (Center for Security and Emerging Technology, December 2023), <https://doi.org/10.51593/2023CA006>; Neil C. Thompson, Kristjan Greenewald, Keeheon Lee, and Gabriel F. Manso, "Deep Learning's Diminishing Returns: The Cost of Improvement Is Becoming Unsustainable," *IEEE Spectrum* 58, no. 10 (October 2021): 50–55, <https://doi.org/10.1109/MSPEC.2021.9563954>.

⁴ Plutarch, "Pyrrhus," in *Plutarch's Lives*, trans. Bernadotte Perrin, vol. 9, Loeb Classical Library 101 (Cambridge, MA: Harvard University Press, 1920), 21.9, https://penelope.uchicago.edu/Thayer/e/roman/texts/plutarch/lives/pyrrhus*.html.

⁵ John von Neumann and Oskar Morgenstern, *Theory of Games and Economic Behavior* (Princeton, NJ: Princeton University Press, 1944), <https://archive.org/details/in.ernet.dli.2015.215284/page/n27/mode/2up>; Thomas C. Schelling, *The Strategy of Conflict* (Cambridge, MA: Harvard University Press, 1980), <https://www.hup.harvard.edu/books/9780674840317>; Robert D. Putnam, "Diplomacy and Domestic Politics: The Logic of Two-Level Games," *International Organization* 42, no. 3 (Summer 1988): 427–60, <https://doi.org/10.1017/S0020818300027697>.

⁶ Kyle Miller and Andrew Lohn, "Onboard AI: Constraints and Limitations" (Center for Security and Emerging Technology, August 2023), <https://doi.org/10.51593/2022CA008>.

⁷ Michael Bowling, Neil Burch, Michael Johanson, and Oskari Tammelin, "Heads-Up Limit Hold'em Poker Is Solved," *Science* 347, no. 6218 (January 9, 2015): 145–49, <https://doi.org/10.1126/science.1259433>.

⁸ Louis Victor Allis, "Searching for Solutions in Games and Artificial Intelligence" (PhD diss., University of Limburg, 1994), <http://fragrieu.free.fr/SearchingForSolutions.pdf>; Jonathan Schaeffer, Neil Burch, et al., "Checkers Is Solved," *Science* 317, no. 5844 (September 14, 2007): 1518–22, <https://doi.org/10.1126/science.1144079>; Bowling et al., "Heads-Up Limit Hold'em Poker Is Solved."

⁹ James Gleick, *Chaos: Making a New Science* (New York: Viking, 1987).

¹⁰ Andrew Lohn, "Hacking AI: A Primer for Policymakers on Machine Learning Cybersecurity" (Center for Security and Emerging Technology, December 2020), <https://doi.org/10.51593/2020CA006>.

¹¹ Christoph Engel, "Dictator Games: A Meta Study," *Experimental Economics* 14, no. 4 (November 2011): 583–610, <https://www.cambridge.org/core/journals/experimental-economics/article/abs/dictator-games-a-meta-study/8597E66280017C3A83863C12A7F8F687>.

¹² Richard H. Thaler, "Anomalies: The Ultimatum Game," *Journal of Economic Perspectives* 2, no. 4 (Fall 1988): 195–206, <https://doi.org/10.1257/jep.2.4.195>.

¹³ Robert W. Rosenthal, "Games of Perfect Information, Predatory Pricing and the Chain-Store Paradox," *Journal of Economic Theory* 25, no. 1 (August 1981): 92–100, [https://doi.org/10.1016/0022-0531\(81\)90018-1](https://doi.org/10.1016/0022-0531(81)90018-1).

¹⁴ Eva M. Krockow, Andrew M. Colman, and Briony D. Pulford, "Cooperation in Repeated Interactions: A Systematic Review of Centipede Game Experiments, 1992–2016," *European Review of Social Psychology* 27, no. 1 (2016): 231–82, <https://doi.org/10.1080/10463283.2016.1249640>; Bela Elmschauser, "Altruism and Ambiguity in the Centipede Game," SocArXiv (2024), <https://doi.org/10.31235/osf.io/93p8s>.

¹⁵ Nicolo Fontana, Francesco Pierri, and Luca Maria Aiello, "Nicer Than Humans: How Do Large Language Models Behave in the Prisoner's Dilemma?" *Proceedings of the International AAAI Conference on Web and Social Media* 19, no. 1 (2025): 522–35, <https://ojs.aaai.org/index.php/ICWSM/article/view/35829>; Kehan Zheng, Jinfeng Zhou, and Hongning Wang, "Beyond Nash Equilibrium: Bounded Rationality of LLMs and

Humans in Strategic Decision-Making," arXiv preprint arXiv:2506.09390 (2025), <https://doi.org/10.48550/arXiv.2506.09390>; Samuel Kirshner, Yiwen Pan, and Jason Xianghua Wu, "Pro-Social When Simple and Cold-Hearted When Complex: How Task Difficulty Shapes LLM Behavior" (SSRN working paper, May 14, 2025), <https://doi.org/10.2139/ssrn.5266963>; Beau G. Schelble, Christopher Flathmann, Lorenzo Barberis Canonico, and Nathan McNeese, "Understanding Human-AI Cooperation Through Game-Theory and Reinforcement Learning Models," in *Proceedings of the 54th Annual Hawaii International Conference on System Sciences* (2021), <https://doi.org/10.24251/HICSS.2021.041>.

¹⁶ Celso M. de Melo, Jonathan Gratch, and Frank Krueger, "Heuristic Thinking and Altruism toward Machines in People Impacted by COVID-19," *iScience* 24, no. 3 (March 2021): 102228, <https://doi.org/10.1016/j.isci.2021.102228>; Mayada Oudah, Kinga Makovi, et al., "Perception of Experience Influences Altruism and Perception of Agency Influences Trust in Human-Machine Interactions," *Scientific Reports* 14 (2024): 12410, <https://doi.org/10.1038/s41598-024-63360-w>.

¹⁷ "AI Is Killing the Web. Can Anything Save It?" *The Economist*, July 14, 2025, <https://www.economist.com/business/2025/07/14/ai-is-killing-the-web-can-anything-save-it>.

¹⁸ James D. Fearon, "Rationalist Explanations for War," *International Organization* 49, no. 3 (Summer 1995): 379–414, <https://doi.org/10.1017/S0020818300033324>.

¹⁹ George A. Akerlof, "The Market for 'Lemons': Quality Uncertainty and the Market Mechanism," *Quarterly Journal of Economics* 84, no. 3 (August 1970): 488–500, <https://www.matetam.com/sites/default/files/themarketforlemons.pdf>.

²⁰ Stuart Armstrong, Nick Bostrom, and Carl Shulman, "Racing to the Precipice: A Model of Artificial Intelligence Development," *AI & Society* 31, no. 2 (2016): 201–6, <https://doi.org/10.1007/s00146-015-0590-y>.

²¹ Barbara W. Tuchman, *The Guns of August* (New York: Ballantine Books, 1994).

²² Fabian Beuke, "Chess Engine Draws," *beuke.org* (blog), October 3, 2025, <https://beuke.org/chess-engine-draws/>.

²³ Kevin Passmore, *The Maginot Line: A New History* (New Haven, CT: Yale University Press, 2025).

²⁴ Ishikawa Group Laboratory, "Janken (Rock-Paper-Scissors) Robot with 100% Winning Rate," Research Institute for Science & Technology, Tokyo University of Science, accessed July 30, 2026, <https://ishikawa-vision.org/fusion/Janken/index-e.html>.

²⁵ Andrew J. Lohn, "Rock, Paper, Scissors, . . . Dynamite - A Model of Disruption from New Technologies," arXiv preprint arXiv:2609.00207 (2026), <https://doi.org/10.48550/arXiv.2609.00207>.

²⁶ Schelling, *The Strategy of Conflict*.

²⁷ Scott E. McIntosh, "The Wingman-Philosopher of MiG Alley: John Boyd and the OODA Loop," *Air Power History* 58, no. 4 (2011): 24–33, <https://www.jstor.org/stable/26276108>.

²⁸ Brian Roberson, "The Colonel Blotto Game," *Economic Theory* 29, no. 1 (2006): 1–24, <https://doi.org/10.1007/s00199-005-0071-5>.

²⁹ Erik Brynjolfsson, Daniel Rock, and Chad Syverson, "Artificial Intelligence and the Modern Productivity Paradox: A Clash of Expectations and Statistics" (NBER Working Paper No. 24001, National Bureau of Economic Research, November 2017), <http://www.nber.org/papers/w24001>; Gene M. Amdahl, "Validity of the Single Processor Approach to Achieving Large Scale Computing Capabilities," *Proceedings of the April 18–20, 1967, Spring Joint Computer Conference* (New York: ACM, 1967): 483–85, <https://www3.cs.stonybrook.edu/~rezaul/Spring-2019/CSE613/reading/Amdahl-1967.pdf>.

³⁰ Edward Geist and Andrew Lohn, "How Might Artificial Intelligence Affect the Risk of Nuclear War?" (RAND Corporation, April 2018), <https://doi.org/10.7249/PE296>.

³¹ Edward Geist, *Deterrence under Uncertainty: Artificial Intelligence and Nuclear Warfare* (New York: Oxford University Press, 2023).

³² Brad Roberts, "On Theories of Victory, Red and Blue" (Livermore Papers on Global Security No. 7, Center for Global Security Research, Lawrence Livermore National Laboratory, June 2020), <https://cgsr.llnl.gov/sites/cgsr/files/2024-08/CGSR-LivermorePaper7.pdf>; Frank G. Hoffman, "The Missing Element in Crafting National Strategy: A Theory of Success," *Joint Force Quarterly* 97 (2020): 55–64,

<https://ndupress.ndu.edu/Media/News/News-Article-View/Article/2106508/the-missing-element-in-crafting-national-strategy-a-theory-of-success/>.

³³ Jacob L. Heim, Zachary Burdette, and Nathan Beauchamp-Mustafaga, "U.S. Military Theories of Victory for a War with the People's Republic of China" (RAND Corporation, February 2024), <https://doi.org/10.7249/PEA1743-1>.

³⁴ Andrew Lohn, "Anticipating AI's Impact on the Cyber Offense-Defense Balance" (Center for Security and Emerging Technology, May 2025), <https://doi.org/10.51593/20250004>; Andrew Lohn, "The Impact of AI on the Cyber Offense-Defense Balance and the Character of Cyber Conflict," arXiv preprint arXiv:2504.13371 (2025), <https://doi.org/10.48550/arXiv.2504.13371>.

³⁵ Kyle Miller, Mia Hoffmann, and Rebecca Gelles, "The Use of Open Models in Research" (Center for Security and Emerging Technology, October 2025), <https://doi.org/10.51593/20240007>.

³⁶ CFR Editors, "U.S.-China Relations," Timelines, Council on Foreign Relations, updated April 15, 2025, <https://www.cfr.org/timeline/us-china-relations>.

³⁷ Carl von Clausewitz, *On War*, trans. J. J. Graham, rev. F. N. Maude (London: Kegan Paul, Trench, Trübner & Co., 1918); Alan Beyerchen, "Clausewitz, Nonlinearity, and the Unpredictability of War," *International Security* 17, no. 3 (Winter 1992–93): 59–90, <https://doi.org/10.2307/2539130>; Defense Technical Information Center, report ADA640735, accessed July 30, 2026, <https://apps.dtic.mil/sti/tr/pdf/ADA640735.pdf>.

³⁸ Mark Rathbone, "Why Were the Vietcong So Hard to Defeat?" *Hindsight* 32, no. 2 (2021–22), <https://magazines.hachettelearning.com/magazine/hindsight/32/2/why-were-the-vietcong-so-hard-to-defeat/>; Benjamin Jensen, "How the Taliban Did It: Inside the 'Operational Art' of Its Military Victory," *Atlantic Council* (blog), August 15, 2021, <https://www.atlanticcouncil.org/blogs/new-atlanticist/how-the-taliban-did-it-inside-the-operational-art-of-its-military-victory/>.

³⁹ Anthony H. Cordesman, "The Key Lessons of America's Recent Wars: Failing or Losing in Grand Strategic Terms" (Center for Strategic and International Studies, June 13, 2023), https://csis-website-prod.s3.amazonaws.com/s3fs-public/2023-06/230613_Cordesman_Lessons_AmericasWar.pdf.

⁴⁰ Andrew Lohn, "Chinese Acquisition of AI Technology," testimony before the U.S.-China Economic and Security Review Commission, April 30, 2026, https://www.uscc.gov/sites/default/files/2026-04/Andrew_Lohn_Testimony.pdf.

⁴¹ Lohn, "Scaling AI."

⁴² Hans Gundlach, Jayson Lynch, and Neil Thompson, "Meek Models Shall Inherit the Earth," arXiv preprint arXiv:2507.07931 (2025), <https://doi.org/10.48550/arXiv.2507.07931>.

⁴³ Gundlach et al., "Meek Models Shall Inherit the Earth."

⁴⁴ Jeffrey Ding, "Running the Right Artificial Intelligence Race: A National Strategy for AI Diffusion" (RAND Corporation, September 2025), <https://doi.org/10.7249/PEA4165-1>.

⁴⁵ Lohn, "Scaling AI."

⁴⁶ Frank Nagle and Daniel Yue, "The Latent Role of Open Models in the AI Economy" (Georgia Tech Scheller College of Business Research Paper No. 5767103, November 18, 2025), <https://doi.org/10.2139/ssrn.5767103>.

⁴⁷ Thomas G. Mahnken, "Cost-Imposing Strategies: A Brief Primer" (Center for a New American Security, November 2014), https://www.files.ethz.ch/isn/186043/CNAS_Maritime4_Mahnken.pdf.

⁴⁸ Anthony Kalashnikov, "Differing Interpretations: Causes of the Collapse of the Soviet Union," *Constellations* 3, no. 1 (2012), <https://doi.org/10.29173/cons16289>.

⁴⁹ Christopher Mims, "Silicon Valley's New Strategy: Move Slow and Build Things," *The Wall Street Journal*, August 1, 2025, <https://www.wsj.com/tech/ai/tech-ai-spending-company-valuations-7b92104b>.

⁵⁰ Arthur C. Clarke, *Profiles of the Future: An Inquiry into the Limits of the Possible* (New York: Harper & Row, 1962).

⁵¹ Miller et al., "The Use of Open Models in Research."

⁵² Harry Hinsley, "The Counterfactual History of No Ultra," *Cryptologia* 20, no. 4 (October 1996): 308–24, <https://doi.org/10.1080/0161-119691884997>; Jack Guy, "Bletchley Park's Contribution to World War II Revealed in New GCHQ Report," CNN,

October 20, 2020, <https://www.cnn.com/2020/10/20/uk/bletchley-park-gchq-wwii-contribution-scli-intl-gbr/index.html>.

⁵³ Kelsey Rightnowar, "Sign Stealing in Baseball Dates Back to the 19th Century. Here's Why the Houston Astros Sign Stealing Was Different," *Frontline*, PBS, October 3, 2023, <https://www.pbs.org/wgbh/frontline/article/sign-stealing-baseball-history/>; David Ubben, "College Football Sign-Stealing: Coaches Weigh in on Michigan, How Common It Is and What Matters," *The Athletic*, October 25, 2023, <https://www.nytimes.com/athletic/4997702/2023/10/25/college-football-sign-stealing-coaches-opinion-michigan/>.

⁵⁴ Auguste Kerckhoffs, "La Cryptographie militaire," *Journal des sciences militaires* 9 (January 1883): 5–38, https://www.petitcolas.net/kerckhoffs/crypto_militaire_1_b.pdf.

⁵⁵ Fearon, "Rationalist Explanations for War."

⁵⁶ Akerlof, "The Market for 'Lemons.'"

⁵⁷ Robert H. B. Netzer and Barton P. Miller, "What Are Race Conditions?" *ACM Letters on Programming Languages and Systems* 1, no. 1 (March 1992): 74–88, <https://doi.org/10.1145/130616.130623>.

⁵⁸ Sven Bostyn and Nicolas Petit, "Patent = Monopoly: A Legal Fiction" (Technology, Innovation and Investment Council, December 2013), https://4ipcouncil.com/application/files/4314/2729/2678/Patent_Monopoly_-_Legal_Fiction_-_Bostyn_and_Petit_-_4iPcouncil.pdf; Zorina Khan, "Are Patents Monopolies?" *Life on the Margin* (blog), July 28, 2021, <https://research.bowdoin.edu/zorina-khan/life-on-the-margin/228/>.

⁵⁹ Bob Metcalfe, "Metcalfe's Law after 40 Years of Ethernet," *Computer* 46, no. 12 (December 2013): 26–31, <https://doi.org/10.1109/MC.2013.374>.

⁶⁰ "How Your Data Is Used to Improve Model Performance," OpenAI Help Center, archived August 16, 2025, <https://web.archive.org/web/20250816225620/https://help.openai.com/en/articles/5722486-how-your-data-is-used-to-improve-model-performance>.

⁶¹ Robin Hanson and Eliezer Yudkowsky, *The Hanson-Yudkowsky AI-Foom Debate* (Berkeley, CA: Machine Intelligence Research Institute, 2013), <https://intelligence.org/files/AIFoomDebate.pdf>.

⁶² Data Is Beautiful, "Most Popular Search Engines 1994–2019," YouTube video, accessed July 30, 2026, <https://youtu.be/yMpj33Y95ZU>; Esteban Ortiz-Ospina, "The Rise of Social Media," Our World in Data, September 18, 2019, <https://ourworldindata.org/rise-of-social-media>; Dominic Rushe, "Myspace Sold for \$35m in Spectacular Fall from \$12bn Heyday," *The Guardian*, June 30, 2011, <https://www.theguardian.com/technology/2011/jun/30/myspace-sold-35-million-news>.

⁶³ Kurt Annen, "On the First Mover Advantage in Stackelberg Quantity Games," *Journal of Economics* 126, no. 3 (April 2019): 249–58, <https://doi.org/10.1007/s00712-018-0622-4>.

⁶⁴ Christopher Cinq-Mars Jarvis, "White Advantage in Chess: An Exploration of Artificial Intelligence," *Cinq-Mars Media* (blog), Medium, July 3, 2017, <https://blog.cinqmarsmedia.com/caissa-b7057651aa01>.

⁶⁵ Beuke, "Chess Engine Draws"; Leonardo Ljubičić (@LeoLjubicic66), X post, August 1, 2020, <https://x.com/LeoLjubicic66/status/1289682018851262465>.

⁶⁶ Fernando F. Suarez and Gianvito Lanzolla, "The Half-Truth of First-Mover Advantage," *Harvard Business Review*, April 2005, <https://hbr.org/2005/04/the-half-truth-of-first-mover-advantage>.

⁶⁷ Tuchman, *The Guns of August*.

⁶⁸ Armstrong et al., "Racing to the Precipice."

⁶⁹ Patrick M. Emerson, "Models of Oligopoly: Cournot, Bertrand, and Stackelberg," chap. 18 in *Intermediate Microeconomics* (Corvallis: Oregon State University, accessed July 30, 2026,) <https://open.oregonstate.edu/intermediatemicroeconomics/chapter/module-18/>.

⁷⁰ Beuke, "Chess Engine Draws."

⁷¹ William E. DePuy, *Selected Papers of General William E. DePuy: First Commander, U.S. Army Training and Doctrine Command, 1 July 1973*, comp. Richard M. Swain, ed.

Donald L. Gilmore and Carolyn D. Conway (Fort Leavenworth, KS: Combat Studies Institute, U.S. Army Command and General Staff College, 1994), 104, <https://www.tradoc.army.mil/wp-content/uploads/2020/10/Selected-Papers-of-General-William-Depuy.pdf#page=104>.

⁷² Ishikawa Group Laboratory, "Janken (Rock-Paper-Scissors) Robot with 100% Winning Rate."

⁷³ Passmore, *The Maginot Line*.

⁷⁴ Yossi Maaravi, Aharon Levy, and Ben Heller, "To Bid or Not to Bid? That Is the Question! First- Versus Second-Mover Advantage in Negotiations," *Negotiation Journal* 39, no. 3 (2023): 259–78, <https://doi.org/10.1111/nejo.12438>.

⁷⁵ Anson Ho, Tamay Besiroglu, et al., "Algorithmic Progress in Language Models," arXiv preprint arXiv:2403.05812 (2024), <https://doi.org/10.48550/arXiv.2403.05812>; Fred (@fredxhans), X post, February 24, 2025, <https://x.com/fredxhans/status/1894148701951611349>.

⁷⁶ John Bansemer and Kyle Miller, "In the News: DeepSeek's Release of an Open-Weight Frontier AI Model" (Center for Security and Emerging Technology, April 22, 2025), <https://cset.georgetown.edu/article/deepseeks-release-of-an-open-weight-frontier-ai-model/>.

⁷⁷ Paul Klemperer, "Auction Theory: A Guide to the Literature," *Journal of Economic Surveys* 13, no. 3 (1999): 227–86, <https://doi.org/10.1111/1467-6419.00083>.

⁷⁸ Sreekaanth S. Isloor and T. Anthony Marsland, "The Deadlock Problem: An Overview," University of Alberta (Wayback Machine, Internet Archive, accessed July 30, 2026), <https://citeseerx.ist.psu.edu/document?repid=rep1&type=pdf&doi=04de4d672e3e3660b53a0e04a43b260235eccc57>.

⁷⁹ "Variations of Rock Paper Scissors," World Rock Paper Scissors Association, archived June 12, 2025, <https://web.archive.org/web/20250612185831/https://wrpsa.com/variations-of-rock-paper-scissors/>.

⁸⁰ Lohn, "Rock, Paper, Scissors, . . .Dynamite."

⁸¹ Dominic Lawson, "Has Chess Got Anything to Do with War?" BBC News, May 3, 2015, <https://www.bbc.com/news/magazine-32542306>.

⁸² Andrew Lohn, "Defending against Intelligent Attackers at Large Scales," arXiv preprint arXiv:2504.18577 (2025), <https://doi.org/10.48550/arXiv.2504.18577>.

⁸³ Lohn, "Rock, Paper, Scissors, . . . Dynamite."

⁸⁴ Maria Dooling, "Develop a Risk Assessment Framework for AI Integration into Nuclear Weapons Command, Control, and Communications Systems" (Federation of American Scientists, June 11, 2025), <https://fas.org/publication/risk-assessment-framework-ai-nuclear-weapons/>.

⁸⁵ Lohn, "Rock, Paper, Scissors, . . . Dynamite."

⁸⁶ Schelling, *The Strategy of Conflict*.

⁸⁷ David E. Hoffman, *The Dead Hand: The Untold Story of the Cold War Arms Race and Its Dangerous Legacy* (New York: Anchor Books, 2010).

⁸⁸ Lohn, "Hacking AI."

⁸⁹ McIntosh, "The Wingman-Philosopher of MiG Alley."

⁹⁰ Roberson, "The Colonel Blotto Game."

⁹¹ John J. Mearsheimer, "Assessing the Conventional Balance: The 3:1 Rule and Its Critics," *International Security* 13, no. 4 (Spring 1989): 54–89, <https://doi.org/10.2307/2538780>.

⁹² Seth G. Jones and Riley McCabe, "Russia's Battlefield Woes in Ukraine" (Center for Strategic and International Studies, June 3, 2025), Figure 6, <https://www.csis.org/analysis/russias-battlefield-woes-ukraine>.

⁹³ G. Ridder, J. S. Cramer, and P. Hopstaken, "Down to Ten: Estimating the Effect of a Red Card in Soccer," *Journal of the American Statistical Association* 89, no. 427 (September 1994): 1124–27, <https://www.jstor.org/stable/2290942>.

⁹⁴ Marco Caliendo and Dubravko Radic, "Ten Do It Better, Do They? An Empirical Analysis of an Old Football Myth" (German Institute for Economic Research, 2006), <https://www.econstor.eu/bitstream/10419/18485/1/dp592.pdf>.

⁹⁵ Cybersecurity and Infrastructure Security Agency, Federal Bureau of Investigation, and Multi-State Information Sharing and Analysis Center, *Understanding and Responding to Distributed Denial-of-Service Attacks* (Washington, D.C.: Cybersecurity and Infrastructure Security Agency, March 21, 2024), https://www.cisa.gov/sites/default/files/2024-03/understanding-and-responding-to-distributed-denial-of-service-attacks_508c.pdf.

⁹⁶ Lohn, "Defending against Intelligent Attackers at Large Scales."

⁹⁷ "The Internet Is Brawling over Whether 100 Men Could Fight a Gorilla," NBC News, May 1, 2025, <https://www.nbcnews.com/pop-culture/pop-culture-news/100-men-fight-gorilla-viral-question-rcna203764>.

⁹⁸ "Statistics on Labour Productivity," ILOSTAT, International Labour Organization, accessed July 30, 2026, <https://ilostat.ilo.org/topics/labour-productivity/>.

⁹⁹ Daron Acemoglu, "The Simple Macroeconomics of AI," *Economic Policy* 40, no. 121 (January 2025): 13–58, <https://doi.org/10.1093/epolic/eiae042>.

¹⁰⁰ Yingying Lu and Yixiao Zhou, "A Review on the Economics of Artificial Intelligence," *Journal of Economic Surveys* 35, no. 4 (September 2021): 1045–72, <https://doi.org/10.1111/joes.12422>; Mariusz-Jan Radlo and Artur F. Tomeczek, "Factors Influencing Labor Productivity in Modern Economies: A Review and Qualitative Text Analysis" (Warsaw School of Economics, 2022), <https://cor.sgh.waw.pl/bitstream/handle/20.500.12182/1012/Rad%C5%82o%20%26%20Tomeczek%20%282022%29%20Factors%20influencing%20labor%20productivity%20in%20modern%20economies.pdf>.

¹⁰¹ Brynjolfsson et al., "Artificial Intelligence and the Modern Productivity Paradox."

¹⁰² Brynjolfsson et al., "Artificial Intelligence and the Modern Productivity Paradox."

¹⁰³ Amdahl, "Validity of the Single Processor Approach."

¹⁰⁴ Miles Brundage, Shahar Avin, et al., "Toward Trustworthy AI Development: Mechanisms for Supporting Verifiable Claims," arXiv preprint arXiv:2004.07213 (2020), <https://doi.org/10.48550/arXiv.2004.07213>.

¹⁰⁵ Andrew Lohn and Micah Musser, "AI and Compute: How Much Longer Can Computing Power Drive Artificial Intelligence Progress?" (Center for Security and Emerging Technology, January 2022), <https://doi.org/10.51593/2021CA009>; Thompson et al., "Deep Learning's Diminishing Returns"; David Patterson, Joseph Gonzalez, et al., "The Carbon Footprint of Machine Learning Training Will Plateau, Then Shrink," arXiv preprint arXiv:2204.05149 (2022), <https://doi.org/10.48550/arXiv.2204.05149>; Cooper Elsworth, Keguo Huang, et al., "Measuring the Environmental Impact of Delivering AI at Google Scale," arXiv preprint arXiv:2508.15734 (2025), <https://doi.org/10.48550/arXiv.2508.15734>.

¹⁰⁶ Matteo Wong, "Silicon Valley's Trillion-Dollar Leap of Faith," *The Atlantic*, July 29, 2024, <https://www.theatlantic.com/technology/archive/2024/07/ai-companies-unprofitable/679278/>.

¹⁰⁷ Rufus Isaacs, "Differential Games I: Introduction" (RAND Research Memorandum RM-1391-PR, RAND Corporation, 1954), https://www.rand.org/content/dam/rand/pubs/research_memoranda/2008/RM1391.pdf; Isaac E. Weintraub, Meir Pachter, and Eloy Garcia, "An Introduction to Pursuit-Evasion Differential Games," arXiv preprint arXiv:2003.05013 (2020), <https://doi.org/10.48550/arXiv.2003.05013>.

¹⁰⁸ Max Hauptman, "Marines Outwitted an AI Security Camera by Hiding in a Cardboard Box and Pretending to Be Trees," *Task & Purpose*, January 24, 2023, <https://taskandpurpose.com/news/marines-ai-paul-scharre/>.

¹⁰⁹ Jason Healey and Tarang Jain, "Are Cyber Defenders Winning?" *Lawfare*, July 14, 2025, <https://www.lawfaremedia.org/article/are-cyber-defenders-winning>.

¹¹⁰ Stephen Biddle, "Back in the Trenches: Why New Technology Hasn't Revolutionized Warfare in Ukraine," *Foreign Affairs*, August 10, 2023, <https://www.foreignaffairs.com/ukraine/back-trenches-technology-warfare>.

¹¹¹ Geist and Lohn, "How Might Artificial Intelligence Affect the Risk of Nuclear War?"

¹¹² Geist, *Deterrence under Uncertainty*.

- ¹¹³ U.S. Central Command (@CENTCOM), "Update from CENTCOM Commander on Operation Epic Fury," X post, March 11, 2026, <https://x.com/CENTCOM/status/2031700131687379148?s=20>; Admiral Brad Cooper, "U.S. Strikes 5,500 Targets in Iran: Operation Epic Fury Update," YouTube video, March 11, 2026, <https://www.youtube.com/watch?v=5L-05hAtkWM>.
- ¹¹⁴ Igor Mikolic-Torreira and Emelia Probasco, "Beyond Targeting: The Untapped Role of AI in Military Decision-Making" (Center for Security and Emerging Technology, July 2026), <https://doi.org/10.51593/20250012>.
- ¹¹⁵ Bonnie Johnson, John M. Green, et al., "Mapping Artificial Intelligence to the Naval Tactical Kill Chain," *Naval Engineers Journal* 135, no. 1 (2023): 155–66, https://nps.edu/documents/10180/142489929/NEJ+Hybrid+Force+Issue_Mapping+AI+to+The+Naval+Kill+Chain.pdf.
- ¹¹⁶ Paul Scharre, *Army of None: Autonomous Weapons and the Future of War* (New York: W. W. Norton, 2018).
- ¹¹⁷ "War as Politics by Other Means," Online Library of Liberty, accessed July 30, 2026, <https://oll.libertyfund.org/pages/clausewitz-war-as-politics-by-other-means>.
- ¹¹⁸ Herbert Lin, "AI in the Information Ecosystem and Its Impact on Nuclear Escalation," *Bulletin of the Atomic Scientists*, March 12, 2026, <https://thebulletin.org/premium/2026-03/ai-in-the-information-ecosystem-and-its-impact-on-nuclear-escalation/>.
- ¹¹⁹ T. T. Parish, "USAMMDA Team Equips a Worldwide Force," U.S. Army, May 2, 2023, https://www.army.mil/article/266314/usammda_team_equips_a_worldwide_force.
- ¹²⁰ Mikolic-Torreira and Probasco, "Beyond Targeting."
- ¹²¹ Forrest E. Morgan, Benjamin Boudreaux, et al., "Military Applications of Artificial Intelligence: Ethical Concerns in an Uncertain World" (RAND Corporation, April 2020), <https://doi.org/10.7249/RR3139-1>.
- ¹²² Lohn, "Hacking AI."
- ¹²³ Siobhán O'Grady and Liubov Sholudko, "Ukraine Remade Air Defense, but Russia Has Changed Its Attacks," *The New York Times*, July 6, 2026,

<https://www.nytimes.com/2026/07/06/world/europe/ukraine-russia-patriot-air-defense.html>.

¹²⁴ Andrew Lohn, "Poison in the Well: Securing the Shared Resources of Machine Learning" (Center for Security and Emerging Technology, June 2021), <https://doi.org/10.51593/2020CA013>.

¹²⁵ Jacob Poushter, Moira Fagan, and Manolo Corichi, "How People around the World View AI" (Pew Research Center, October 15, 2025), <https://www.pewresearch.org/global/2025/10/15/how-people-around-the-world-view-ai/>.

¹²⁶ Phil Stewart and Idrees Ali, "Lawmakers Demand Pentagon Release Findings from Probe of Iran School Strike," Reuters, July 13, 2026, <https://www.reuters.com/legal/government/lawmakers-demand-pentagon-release-findings-probe-iran-school-strike-2026-07-13/>.

¹²⁷ Morgan, Boudreaux, et al., "Military Applications of Artificial Intelligence."

¹²⁸ R. Harrison Wagner, "Bargaining and War," *American Journal of Political Science* 44, no. 3 (July 2000): 469–84, <https://doi.org/10.2307/2669259>.

¹²⁹ Michael C. Horowitz and Lauren Kahn, "First Ukraine, Now Iran: A New Era of Drone Warfare Takes Hold," Council on Foreign Relations, March 9, 2026, <https://www.cfr.org/articles/the-new-era-of-drone-warfare-takes-root-in-iran>.