



NATIONAL DEFENSE RESEARCH INSTITUTE

CHILDREN AND FAMILIES
EDUCATION AND THE ARTS
ENERGY AND ENVIRONMENT
HEALTH AND HEALTH CARE
INFRASTRUCTURE AND
TRANSPORTATION
INTERNATIONAL AFFAIRS
LAW AND BUSINESS
NATIONAL SECURITY
POPULATION AND AGING
PUBLIC SAFETY
SCIENCE AND TECHNOLOGY
TERRORISM AND
HOMELAND SECURITY

The RAND Corporation is a nonprofit institution that helps improve policy and decisionmaking through research and analysis.

This electronic document was made available from www.rand.org as a public service of the RAND Corporation.

Skip all front matter: [Jump to Page 1](#) ▼

Support RAND

[Browse Reports & Bookstore](#)

[Make a charitable contribution](#)

For More Information

Visit RAND at www.rand.org

Explore the [RAND National Defense Research Institute](#)

View [document details](#)

Limited Electronic Distribution Rights

This document and trademark(s) contained herein are protected by law as indicated in a notice appearing later in this work. This electronic representation of RAND intellectual property is provided for non-commercial use only. Unauthorized posting of RAND electronic documents to a non-RAND website is prohibited. RAND electronic documents are protected under copyright law. Permission is required from RAND to reproduce, or reuse in another form, any of our research documents for commercial use. For information on reprint and linking permissions, please see [RAND Permissions](#).

This report is part of the RAND Corporation research report series. RAND reports present research findings and objective analysis that address the challenges facing the public and private sectors. All RAND reports undergo rigorous peer review to ensure high standards for research quality and objectivity.



NATIONAL DEFENSE RESEARCH INSTITUTE

Brandishing Cyberattack Capabilities

Martin C. Libicki

Prepared for the Office of the Secretary of Defense

Approved for public release; distribution unlimited

The research described in this report was prepared for the Office of the Secretary of Defense (OSD). The research was conducted within the RAND National Defense Research Institute, a federally funded research and development center sponsored by OSD, the Joint Staff, the Unified Combatant Commands, the Navy, the Marine Corps, the defense agencies, and the defense Intelligence Community under Contract W74V8H-06-C-0002.

The RAND Corporation is a nonprofit institution that helps improve policy and decisionmaking through research and analysis. RAND's publications do not necessarily reflect the opinions of its research clients and sponsors.

RAND® is a registered trademark.

© Copyright 2013 RAND Corporation

Permission is given to duplicate this document for personal use only, as long as it is unaltered and complete. Copies may not be duplicated for commercial purposes. Unauthorized posting of RAND documents to a non-RAND website is prohibited. RAND documents are protected under copyright law. For information on reprint and linking permissions, please visit the RAND permissions page (<http://www.rand.org/publications/permissions.html>).

Published 2013 by the RAND Corporation
1776 Main Street, P.O. Box 2138, Santa Monica, CA 90407-2138
1200 South Hayes Street, Arlington, VA 22202-5050
4570 Fifth Avenue, Suite 600, Pittsburgh, PA 15213-2665
RAND URL: <http://www.rand.org>
To order RAND documents or to obtain additional information, contact
Distribution Services: Telephone: (310) 451-7002;
Fax: (310) 451-6915; Email: order@rand.org

Preface

The U.S. military exists not just to fight and win wars but also to deter them and even dissuade others from preparing for them. Deterrence is possible only when others have a good idea of what the U.S. military can do. Such acknowledgment is at the heart of U.S. nuclear deterrence strategy and, to a lesser extent, our maintaining strong mobile conventional forces that can intervene almost anywhere on the globe. Cyberattack capabilities, however, resist such demonstration, for many reasons, not least of which is that their effects are very specific to details of a target system's software, architecture, and management. But the fact that cyberattack capabilities cannot easily be used to shape the behavior of others does not mean they cannot be used at all. This report explores ways that cyberattack capabilities can be "brandished." It then goes on to examine the obstacles to doing so and sketches some realistic limits on our expectations.

This research was sponsored by the Office of the Secretary of Defense and conducted within the International Security and Defense Policy Center of the RAND National Defense Research Institute, a federally funded research and development center sponsored by the Office of the Secretary of Defense, the Joint Staff, the Unified Combatant Commands, the Navy, the Marine Corps, the defense agencies, and the defense Intelligence Community.

For more information on the RAND International Security and Defense Policy Center, see <http://www.rand.org/nsrd/ndri/centers/isdp.html> or contact the director (contact information is provided on the web page).

Contents

Preface	iii
Summary	vii
Acknowledgments	xiii
 CHAPTER ONE	
No May Day Parades	1
Background and Purpose	1
What Is Brandishing?	2
Brandishing and Deterrence: A Cautionary Note	3
Organization of This Report	4
 CHAPTER TWO	
The Broad Effects of Brandishing Cyberattack Capabilities	5
What Role for Brandishing?	5
Would a Successful Penetration Say Enough About What Cyberwar Can Do?	6
Inducing Fear, Uncertainty, and Doubt	8
Would Such a Strategy Work with Russia and China?	10
How the Fear of Penetration Might Affect Enemy Operational Behavior	10
How Fears of Penetration Might Affect Defense Investments	12
The Algebra of Direct Intimidation	13
Paradoxes of Intimidation	16
U.S. Policy and the Legitimization of Cyberwar	16
 CHAPTER THREE	
Brandishing Cyberattack in a Nuclear Confrontation	19
Two-Party Confrontations	20
Disabling a Capability Versus Thwarting a Threat	23
The Rogue State Might Try to Discredit the Cyberwar Bluff	23
Can Cyberattack Brandishing Forestall Unilateral Nuclear Use or Threat of Use?	25
Friendly Third Parties Add Complications	26
Summation	27
 CHAPTER FOUR	
Conclusions	29
References	31

Summary

Background and Purpose

The U.S. military exists not just to fight and win wars but also to deter them, that is, to persuade others not to start them (or even prepare for them). Deterrence is possible only when others know or at least have good indications of what the U.S. military can do. Such acknowledgment is at the heart of U.S. nuclear deterrence strategy and, to a lesser extent, the U.S. maintaining strong mobile conventional forces that can intervene almost anywhere on the globe.

Cyberattack capabilities resist such demonstration. No one knows exactly or even approximately what would happen if a country suffered a full-fledged cyberattack, despite the plethora of hostile activity in cyberspace. For one thing, there has never been a cyberwar—attacks with destruction and casualties comparable to physical war. Theory also works against demonstration. Flaws in target systems enable cyberattacks. To reveal which flaws enable attack is to inform others how to fix the flaws and hence neutralize them. It is no wonder that national cyberwar capabilities are a closely guarded secret.

That cyberattack capabilities cannot *easily* be used to shape the behavior of others does not mean they cannot be used at all. This report explores ways that cyberattack capabilities can be “brandished” and the circumstances under which some deterrence effect can be achieved.¹ It then goes on to examine the obstacles to realizing such achievement and sketches out some realistic limits on the expectations.

As a matter of policy, the United States has never said that it would use cyberattacks, but neither has it said that it would not. It has also not vigorously disputed the notion that it had some hand in the Stuxnet attacks on the Iranian nuclear facility.

The Broad Effects of Brandishing Cyber Capabilities

Any state that would discourage other states from aggression in the physical or cyber world by brandishing cyberattack capabilities should first ask itself whether the point of doing so is to look powerful or to make others look powerless. Although both aims are useful, the need to concentrate on one message in a strategic communications campaign suggests the usefulness of making a choice. Emphasizing one’s power has the advantage of inducing caution in all actual or potential opponents and deflects predators to easier prey. It may also reflect well on other

¹ Note that the usage of *brandishing* here is intended to invoke the imagery of warriors displaying their weapons (and hence their capabilities) before battle, by way of warning, rather than that of a criminal displaying a gun to threaten a victim.

sources of national power. But trumpeting the weaknesses of others deters troublesome states by reminding them of their vulnerabilities. It also deflects the accusations of self-promotion by turning the focus toward potential victims.

A bigger challenge is *how* to demonstrate cyberwar capabilities. The most obvious way to demonstrate the ability to hack into an enemy's system is to actually do it, leave a calling card, and hope it is passed forward to national decisionmakers. If the attack can be repeated at will or if the penetration can be made persistent, the target will be forced to believe in the attacker's ability to pop into his system at any time. This should force the target to recalculate its correlation of forces against the attacker.

But as with many things in cyberspace, it sounds simpler than it is. Hinting at outright success is difficult without conceding one's participation in mischief in the first place and hence cyberwar's legitimacy as a tool of statecraft, something countries only started acknowledging in mid-2012. Targets of little value tend to be easy, but penetrating them is unimpressive. Targets of some value are, for that reason, much harder, often because they are electronically isolated. Finally, the ability to penetrate a system does not necessarily prove the ability to *break* a system. The latter requires not only breaking into sufficiently privileged levels but also figuring out how to induce a system to fail *and keep on failing*. But penetration may be sufficiently scary in itself if the target leadership cannot discern the difference between breaking into and breaking.

Breaking a system is more hostile and more difficult than breaking into one. It requires an understanding of what makes the system fail. Getting the desired results also requires shaping the attack so that those who administer the system cannot detect the attack and repair the damage quickly. Conveying to others the ability to bring their systems down and keep them down is not easy. Intended audiences of such demonstrations may subsequently identify the flaw that would allow such an attack and fix it. If so, for brandishing to work, cyberattack capabilities may require repeated demonstration. Alternatively, a less hostile demonstration could be to manipulate the system but not to the point of harming it, a fine line.

Can brandishing help dissuade other states from pursuing a network-centric high-technology force to counter U.S. military capabilities? The best way to demonstrate the risk of network-centricity is to hack into military systems to show their fragility (claiming responsibility is unnecessary; the point is to emphasize not U.S. power but the vulnerability of the enemy's network-centric systems). In other circumstances, making what is vulnerable clear may be unnecessary, perhaps unwise. Every hack leads to fixes that make the next exploitation much harder. But the hint of an attack that leaves no specific trace leaves nothing specific to fix. The point is to convince others that they cannot protect their systems even after paying close attention to their security. The vulnerability of less sophisticated states to unseen manipulation may be higher when the target does not really understand the technology behind its own weapon systems. Often, the target's lack of access to others' source code and not having built any of its own complicates figuring out what went wrong and how to fix it.

Not all states will throw up their hands, though. Some may reason that, because the effects of cyberattacks are temporary and difficult, their systems can survive the initial exchange and recover for subsequent rounds. So, they pursue high technology and ignore the demonstrated possibility that high-technology military campaigns might last days rather than months or years. A subtler counterstrategy is to network warfighting machines (configured not to touch the Internet) and forget about networking people; isolation avoids some of the pesky vulnerabilities arising from human error (notably those associated with authentication, such as pass-

words and tokens). Or they simply renounce network-centric warfare and conclude that they avoided the pitfalls of depending on technology.

It is unclear whether brandishing cyberattack capabilities can curb the enthusiasm of potential foes for war. Some states may feel they have little choice. Others may feel that they can succeed even if their high-technology systems fail. Yet others may discount the possibility entirely, believing their systems—when called on for war—would be disconnected from the rest of the world. Last, the target may simply not believe its own vulnerability, not during peacetime and certainly not when the war drums sound. Going to war requires surmounting a great many fears; digital ghosts may simply be another.

The unwanted effects of making even some third parties believe that we have invaded their systems warrants note. *All* other militaries may also shy away from foreign sources for logic-processing devices (whether software or hardware) and may redouble their efforts to increase their indigenous production capability or, alternatively, pressure their suppliers to hand over source code with their systems, a negative if their supplier is a U.S. corporation. The problem does not go away if the threat turns out not to work. Countries certain that their military systems have been invaded may blame the United States for any military failures even with no evidence of U.S. involvement. Conversely, the United States may be accused of complicity with a rogue state whenever its equipment does *not* fail because this could only mean that the United States condoned the rogue's actions.

Brandishing Cyberattack Capabilities in a Nuclear Confrontation

Are there circumstances in which the United States might usefully hint that it could interfere with a rogue state's nuclear weapons and thereby defuse a nuclear confrontation? Posit a rogue state with dozens of weapons capable of hurting neighbors but not the United States. Assume further the United States has a robust cyberwar capability from which the rogue state's nuclear arsenal is not provably immune. To the extent that the rogue state is far more willing to go to the brink than the United States is, it may not be completely deterred by the U.S. promise of a devastating reaction to its nuclear use. The rogue nuclear state, we further posit, threatens that, if the United States crosses its "red line," it could or would respond with a nuclear shot.

We first model a two-state confrontation and then introduce a friendly third state on whose behalf the United States is acting.

The question is, which is more implacable: the United States determined to cross the red line or the rogue state equally determined to respond with nuclear weapons? If one side can communicate enough confidence in its willingness to keep pressing, the other side may feel that the first side will not back down and would thus logically recognize that the choice is between yielding and catastrophe. The more that the other side indicates it might yield, the greater the impetus for the first side to stand firm, making it seem even more implacable to the other side.

The purpose of brandishing a cyberwar weapon is to threaten the other side's ability use its nuclear capability in a crisis. This purpose is less to make the other side doubt its own nuclear capability—although that can help—but to project a belief that the United States will press on either because the rogue state's weapons will not work or because the rogue state will respond to the brandisher's confidence (underwritten, of course, by its deterrence capability) and back down. Note that the logic works even if the target state believes that the brandisher's

confidence has no basis in reality (i.e., its own nuclear command and control is rock solid). The rogue state needs only to believe that the *brandisher* believes it can act with impunity to conclude that the choice is between disaster and backing down. To be sure, because a cyberwar capability cannot be tested in the same way that an antimissile capability can be tested, the rogue state may conclude that the brandisher's confidence is unwarranted and therefore that such confidence should not exist and hence *does not* exist. But that could also be wishful thinking on the rogue state's part.

If brandishing a cyberthreat created a use-it-or-lose-it dilemma for the rogue state leading to nuclear use, brandishing could backfire on the United States. But it should not, largely because it is not a threat of what *will* happen but what has *already* happened: The flaw has already been exploited. However, brandishing a cyberwar capability, particularly if specific, makes it harder to *use* such a capability because brandishing is likely to persuade the target to redouble its efforts either to find or route around the exploited flaw (the one that enabled the United States to neutralize its nuclear threat). Brandishing capabilities sacrifices the ability to manage a war in exchange for the ability to manage a crisis.

One possible component of the brandishing process is to convey that a nuclear shot that failed will be noticed—and responded to—even if the failure would be invisible to outside observers. Otherwise, the rogue state may reason that failure is costless and that success, while potentially very costly, at least demonstrates that the rogue state is serious. But if the induced failure is not obvious (e.g., the button is pushed and nothing happens), can the United States retaliate against an attempted action that only the United States saw?

Once third parties are in a position to veto U.S. military actions, they can complicate the use of brandishing. Although third parties may have greater animus against the nuclear-armed state and, correspondingly, a greater willingness to see it humiliated, and certainly deterred, they may well blanch at the cyberwar-backed bluff. First, they and their citizens are likely to be at greater risk by dint of sitting within range of the rogue state's nuclear weapons. Second, they would know little about U.S. cyberwar capabilities and may thus have less confidence that such capabilities would work than the United States (supposedly) has. The rogue state may figure that it need not stare down the United States if it can scare the third party whose concurrence is needed for U.S. actions.

The United States may need options to convince the third party that it can stand fast because, among other things, its cyberwar capabilities will neutralize the nuclear threat. It could say, "trust me on this" or else. But a U.S. response that goes beyond asking for trust may have to reveal much more about the details of U.S. cyberwar capabilities than the United States seems comfortable doing today. A crisis makes revelation problematic: Even though steadfastness requires pro-U.S. forces to project faith in the U.S. ability to nullify a nuclear threat, those nervous of taking such a huge risk, skeptics of cyberwar's power, or opponents of the United States within the government have every incentive to cast doubt on the proposition or even leak the information entrusted to them. (Incidentally, a similar logic applies if the friendly third party is domestic, such as the U.S. Congress, opinion makers, and the public.) It may be to the rogue state's advantage to imply that cyberwar capabilities (rather than the confidence in the deterrence effect of its nuclear weapons) are the *primary* basis for the firm stance the United States has adopted. This could pressure the United States to demonstrate what it can do.

Conclusions

Brandishing a cyberattack capability would do three things: declare a capability, suggest the possibility of its use in a particular circumstance, and indicate that such use would really hurt. In the era of the U.S.-Soviet nuclear standoff, the suggestion of use was the most relevant. Possession was obvious, and its consequences were well understood. The same does not hold true for cyberweapons. Possession is likely not obvious, and the ability to inflict serious harm is debatable. Even if demonstrated, what worked yesterday may not work today. But difficult does not mean impossible.

Advertising cyberwar capabilities may be helpful. It may back up a deterrence strategy. It might dissuade other states from conventional mischief or even from investing in mischief-making capabilities. It may reduce the other side's confidence in the reliability of its information, command and control, or weapon systems. In a nuclear confrontation, it may help build the edge that persuades other states that the brandisher will stay the course, thereby persuading them to yield.

Yet proving such capability is not easy, even if it exists. Cyber capabilities exist only in relationship to a specific target, which must be scoped to be understood. Cyber warriors can illustrate their ability to penetrate systems, but penetration is not the same as getting them to fail in useful ways. Since cyberattacks are essentially single-use weapons, they are diminished in the showing. It can be hard to persuade your friends that you have such capabilities when skepticism is in their interest.

Furthermore, brandishing may backfire. Touting an ability to strike back in cyberspace may communicate a tendency to shy from violence. Claiming the power to alter reality may convince others to blame the claimant when reality is disagreeable. Interfering with others' command and control may allow them to justify rules of engagement that abdicate their own responsibility over subordinates. And asserting an ability to nullify opposing nuclear systems may spur them to call what they perceive as a bluff.

Should the United States put the world on notice that it has cyber capabilities and knows how to use them? The wisdom of that course is not obvious. Evidence is scant that others act because they do not believe the United States has or can develop cyber capabilities. Conversely, the gains from brandishing such capabilities depend on the context and can be problematic even then.

There is both promise and risk in cyber brandishing, in both the conventional and nuclear cases. It would not hurt to give serious thought to ways in which the United States can enhance its ability to leverage what others believe are national capabilities. Stuxnet has certainly convinced many others that the United States can do many sophisticated things in cyberspace (regardless of what, if anything, the United States actually contributed to Stuxnet). This effort will take considerable analysis and imagination, inasmuch as none of the various options presented here are obvious winners. That said, brandishing is an option that may also not work. It is no panacea, and it is unlikely to make a deterrence posture succeed if the other elements of deterrence (e.g., the will to wage war or, for red lines drawn in cyberspace, the ability to attribute) are weak.

Acknowledgments

The author would like to acknowledge the valuable contribution to the analysis contained in this report made by Roger C. Molander, who passed away on March 25, 2012. Dr. Molander spent a full career working to analyze and illuminate the complex dimensions of strategic nuclear deterrence. More recently, he had applied the same analytic rigor to questions raised by cyber operations, particularly the interaction of technical and operational factors with political incentives. He made innovative contributions to understanding the difficult questions posed by large nuclear weapons holdings and by powerful cyber operations capabilities. The latter formed an important part of the intellectual foundation for the material presented in this report, notably Chapter Three, which arose from conversations he initiated.

The author also acknowledges the generously provided and very useful commentary from Stuart Johnson, James Dobbins, Forrest Morgan, and the formal reviews of David C. Gompert and James T. Quinlivan. Finally, the U.S. Naval War College sponsored production of an early version of the material in Chapter Two.

No May Day Parades

Background and Purpose

Marching warfighters and weaponry down urban thoroughfares has been a time-honored way for states to hint at their ability to carry out war. Cyberwar capabilities, to be sure, resist such presentation. Cadres of computer geeks advancing with laptops in their rucksacks somehow do not inspire the same awe.

The inability to display power points to a larger dilemma of cyberwar. The U.S. military exists not just to fight and win wars but also to deter them, that is, to persuade others not to start them (or even prepare for them). To do this, it helps to demonstrate that the U.S. military is and always will be likely to ruin those who would fight it—whether the ruin be a crushed military or a damaged society. By so doing, the United States may hope to deter others from attacking it or its vital interests—either kinetically or via cyberspace. It may even hope to dissuade states from developing digitized capabilities that are particularly vulnerable to cyberattack. Although May Day parades are a bit of a caricature, a state would rationally examine the ability of its potential adversaries before pursuing its politicomilitary strategies. But cyberwar capabilities are hard to examine.

Why so? No one doubts what would happen if a nuclear-armed power dropped its big weapon on a city, even though no city has been hit by a nuclear bomb since 1945. The physics are clear, and they work anywhere. But no one knows exactly or even approximately what would happen if a country suffered a full-fledged cyberattack, despite the plethora of hostile activity in cyberspace that shows no signs of abating. For one thing, there has never been such an attack.

Theory also discourages good a priori expectations. First, systems are vulnerable only to the extent that they have exploitable errors that their owners do not know about or have simply ignored. Second, even if a cyberattack works, the damage it wreaks tends to be proportional to the time required to recover the attacked system, something neither the defender nor the attacker can easily predict. Third, national cyberwar capabilities are a closely guarded secret.

Having spent much time and trouble developing cyberwar capabilities, states thus have nothing to show for their efforts until and unless they go to cyberwar. Although some of the capabilities needed for cyberwar are the same ones used for cyberespionage, some are not. Bringing systems down requires effort to understand their failure modes; keeping them down requires being able to insert code into the target networks and system in ways that make it difficult to eradicate. Furthermore, systems targeted by espionage (e.g., email networks) are very different from the harder systems that run critical infrastructure or war machines.

That cyberattack capabilities cannot *easily* and credibly be brandished does not mean they cannot be brandished at all. This report explores ways that cyberwar capabilities can be so used, obstacles to doing so well, some uses of doing so, some risks involved, and limits on our expectations.

What Is Brandishing?

Brandishing a weapon communicates what it is and suggests how it would be used.¹ Brandishing can be implicit, leaving it to others to determine the implications of its use. Or it can be explicit, with the owner choosing the context and timing for signaling something.²

Capabilities are generally brandished to shape or at least reinforce other states' estimates of the risks they face by opposing the brandisher. For cyberspace, estimates vary greatly. Cyberattack capabilities are always capabilities against specific systems, and states vary in what systems they have, how important they are, and how secure they are.

Because no state where news about Stuxnet has penetrated can seriously believe the United States lacks offensive cyberattack capabilities and because so many argue for the primacy of offense therein,³ U.S. cyberwar capabilities may already be discouraging others from mischief today.⁴ Weapons alone can do this. In 1932 (before Germany had a Luftwaffe), Stanley Baldwin persuaded the British Parliament not to intervene too hastily in European affairs by arguing that a serious adversary could use airpower to do great damage to Great Britain: "The bomber will always get through."⁵

Why brandish cyberattack capabilities at all?

- One reason is simply to make a threat, either specifically ("do this and we will carry out a cyberattack") or generally ("do this, and we will respond—with capabilities that include a possible cyberattack").

¹ Note that the usage of *brandishing* here is intended to invoke the imagery of warriors displaying their weapons (and hence their capabilities) before battle, by way of warning, rather than that of a criminal displaying a gun to threaten a victim.

² The explicitness of the threat does not necessarily conform to how openly a capability is declared. It is possible to be very open about having a capability without drawing red lines. (A red line is a limit a state establishes beyond which it feels obliged to take action.) With somewhat more difficulty, one can make an explicit threat based on a coyly presented capability.

³ Among the many sources that argue that offense is dominant in cyberspace are Jonathan Masters, "Confronting the Cyber Threat," New York: Council on Foreign Relations, May 23, 2011; Richard J. Harknett, John P. Callaghan, and Rudi Kauffman, "Leaving Deterrence Behind: War-Fighting and National Cybersecurity," *Journal of Homeland Security and Emergency Management*, Vol. 7, No. 1, November 11, 2010; and Eric Sterner, "Stuxnet and the Pentagon's Cyber Strategy," Arlington, Va.: George C. Marshall Institute, October 13, 2010.

⁴ Nevertheless, when asked whether the United States had ever "'demonstrated capabilities' in cyberspace in a way that would lead to deterrence of potential adversaries," General Alexander responded, "Not in any significant way." Keith Alexander, "Advance Questions for Lieutenant General Keith Alexander, USA, Nominee for Commander, United States Cyber Command," statement to the U.S. Senate Committee on Armed Services, April 15, 2010, p. 21.

⁵ George H. Quester, *Deterrence Before Hiroshima*, Piscataway, N.J.: Transaction Publishers, 1986. Note that Baldwin was speaking over a dozen years and many generations of aircraft after the last use of airpower against a sophisticated foe. Yet as the Battle of Britain later proved, once countries faced real bombers, damage was less than feared, and they did not always get through.

- Another is to counter a threat, whether explicit or implicit. This is similar to announcing a capability for ballistic missile defense after the other side has announced a ballistic missile capability—with cyberwar playing the role of a weapon aimed at the missile’s command and control. Such an announcement may be made to downplay the threat, assuring oneself and allies and thereby weakening the threat’s deterrent power. If the underlying threat is itself a counterdeterrent (“if you launch a missile, we will launch one back”), the cyberattack capability can be brandished to reinforce the original deterrent (“yes, but your missile will fail, and so we will ignore your threat”). Such brandishing helps project confidence.
- Brandishing a cyberattack capability can warn others against pursuing capabilities that depend on digital systems in general and networks in particular. A variant of that threat is to hint that the information that potential foes use to make operational or even strategic decisions may be corrupted and is therefore unreliable. The threat need not be proactive (“if you do this . . . ”); the brandisher can hint that a corruption attack has already reached its target, meaning that even existing data cannot be trusted.

The credibility of the cyberattack threat will depend on a state’s track record in cyberspace coupled with its general reputation at military technology and the likelihood that it would use such capabilities when called on. Finally, as the technologies of cyberspace and the targeted state’s dependence on cyberspace evolve, so too will the effectiveness of such threats.

Brandishing and Deterrence: A Cautionary Note

One reason for a state to say or hint that it has offensive cyberwar capabilities is to give teeth to a deterrence policy.⁶ As a general rule, the greater a state’s capabilities to strike, the greater the consequences of other states of crossing the lines it lays down, and thus the lower the likelihood that other states will cross the lines (at least up to the point at which other states fear for their sovereignty and try to cut the state down to size because it is so threatening). That noted, deterrence also requires some clarity on where the red lines are and how willing such a state is to carry out its threat and by what means. Absent such clarity, brandishing may have an effect opposite from the one intended.

Much depends, therefore, on what other states conclude about the motive for brandishing a cyberwar capability and the timing of the brandishing. If the threatening state is explicit that it will use cyber means to retaliate for crossing certain red lines (presumably, but not necessarily, in cyberspace), the role of brandishing is fairly clear: to give substance to a threat. But the timing may raise questions, especially if other states do not learn anything new about the threatening state’s capabilities (which they always assumed existed) but were uncertain about why the threatening state believed the point had to be made explicit. Context would matter. Brandishing a capability to reinforce a threat that has just been made (or a red line that has just been laid down or redrawn) may raise a few questions of timing, but brandishing a capability out of the blue might raise more. Some may view it as a bluff, an attempt to put a brave face

⁶ Consistent with the author’s previous report on deterrence (Martin C. Libicki, *Cyberdeterrence and Cyberwar*, Santa Monica, Calif.: RAND Corporation, MG-877-AF, 2009), the word *deterrence* refers only to deterrence by punishment and does not include deterrence by denial.

on the discovery that cyber capabilities are not impressing others for the good reason that they are *not* all that impressive.

If the threatening state, however, has not stated or strongly hinted that its choice of retaliatory weapon sits in cyberspace, other states may wonder why it is emphasizing its retaliatory capabilities *in that domain*. True, the answer may be innocent: A bureaucratic struggle may have been resolved, or a new cyberattack capability may be deemed mature. But states not privy to such explanations may conclude that brandishing a cyberattack capability was meant to signal that more violent responses are off the table. States that do not fear cyber capabilities (maybe because their militaries or economies are not all that digitized) may therefore conclude that they can relax and may therefore be *less* deterred.

Organization of This Report

With these cautions out of the way, the remainder of the report examines the consequences of brandishing cyberattack capabilities. I examine four separate goals for brandishing cyberattack capabilities: to discourage military operations; to dissuade countries from investing in network capabilities; to permit the United States to face down nuclear-armed rogue states; and to inhibit unprovoked nuclear aggression.

Chapter Two is a general treatment of brandishing: whether and how states can prove or at least back up claims that they have such a capability and against whom, how it might be used to reduce the desire of other states to carry out operations or even invest in certain operational capabilities, and the calculus and paradox of intimidation.

Chapter Three specifically treats how cyberweapons may be brandished in a nuclear confrontation. Clearly, when facing obliteration, the threat of being hacked is unlikely to register very high. However, the operational use of cyberwar to thwart an opponent's nuclear command-and-control cycle may play a more interesting role.

Chapter Four wraps up the key insights.

The Broad Effects of Brandishing Cyberattack Capabilities

Brandishing a capability that cannot be displayed for inspection and cannot be demonstrated in any detail without rapidly nullifying it is more than a little challenging. In this chapter, I examine various ways of addressing the challenge, concluding that, while each has its merits, none is altogether satisfactory. In sequence, therefore, the chapter discusses how system penetration may allude to cyberattack capabilities, how the fear that penetration has already occurred may be created and sustained, and how fears of penetration may effect an adversary's operational behavior or even its defense investments. It then examines some consequences of employing cyberattacks as a coercive device, discusses ways in which brandishing may backfire, and concludes by touching on current policies associated with the legitimization of cyberwar.¹

What Role for Brandishing?

Because the potential for cyberattacks arises from the target's vulnerabilities *coupled* with the attacker's ability to exploit them, is the desired effect of brandishing cyberattack capabilities to look powerful or to make the other side look powerless? Of course, the answer could be both, and both may be useful, but if the brandishing is part of an overall strategic communications campaign, it may help to decide what to emphasize in such a campaign.

Looking powerful is the more efficient option. It induces caution in actual or potential opponents. The demonstration does not have to be repeated for each one. Looking large also serves to deflect potential attackers away from one state toward others. Finally, there is glory in it; success reflects well on other sources of national power.

But concentrating instead on exposing another state's weaknesses also has its virtues. It serves to deter troublesome states by reminding them of their vulnerabilities. It also deflects the accusations of self-promotion ("look at how powerful I am") by turning the focus toward others. After all, a state shown to be vulnerable to one attacker in cyberspace may be presumed vulnerable to others. Even if the state retaliates, its systems will still be vulnerable and perceived as such.

For the United States, a further goal may perhaps be to discourage attacks on *anyone*. In a globalized economy, a severe cyberattack against foreign institutions may hurt the United States in its pocketbook: directly, if the U.S. economy relies on their information services, and

¹ An early version of the core argument of this chapter can be found in the author's "Wringing Deterrence from Cyberwar Capabilities," in Richmond M. Lloyd, ed., *Economics and Security: Resourcing National Priorities*, proceedings of a workshop sponsored by the William B. Ruger Chair of National Security Economics, Newport, R.I.: Naval War College, May 19–21, 2010, pp. 259–272.

indirectly, through the effects on export prices and availability. Such an attack may have political ramifications. Cyberwar erodes trust; successful attacks confound the rule of law. A posture to inhibit cyberwar in general, rather than just on the United States, fits with the current U.S. policy narrative that today's security problems are the results of rogue action by rogue states.

Would a Successful Penetration Say Enough About What Cyberwar Can Do?

The most obvious way of demonstrating the ability to hack into someone else's system is to actually do it and leave a calling card (e.g., "Kilroy was here"). The effect need not be obvious to the public, but it must at least be obvious to system administrators. If the attack can be repeated at will or if the penetration can be made ineradicable, the target may be forced to believe that the perpetrator's ability to pop into the target's system at will is a fact. This forces the target to recalculate its correlation of forces against the perpetrator.

This sounds simple. As with most things in cyberspace, it is not. The first problem is whether the calling card would be read and its existence transmitted to the leadership. If it is simply left for someone to stumble over, the answer may be "no." Ironically, the more penetrable the system is, the less astute its administrators are, all else being equal. Thus, the odds of having the penetration discovered and transmitted up the line go down. For this reason, any calling card may have to be more obvious. Perhaps it can email itself, so to speak, to the system's administrators in the hopes that they will tell the leadership. If the target system is connected to the rest of world—a big *if* for sensitive systems—it can email itself directly to the target's leadership. That should work (unless the leaders get it into their heads that it was a trick by their own cyberwar proponents to gain more resources for cyber security). The opposite is also possible. If acknowledging a penetration is embarrassing and puts jobs and, in some countries, lives at risk, such hints may be erased by embarrassed system administrators. Revealing a secret that could only have been stolen from such a system eliminates the problem of post hoc erasure but introduces the question of whether the information could have come only from system penetration (rather than, say, spies).

The next difficulty is proving that the ability to penetrate a system at will is something that matters. If, as noted, a proven penetration is a one-time event, the target may convince itself that it can take measures to ensure that a repeat performance will be impossible. Or the victim may tolerate the attacker's ability to stay on the system precisely because it finds penetration less intolerable than the cost of hitting a systemic reset button. Such an assessment automatically puts an upper limit on the demonstration effect of the cyberattack. Furthermore, the effectiveness of the penetration has everything to do with the sensitivity of the system being penetrated. This requires understanding which systems are critical to the target and whose penetration would be impressive. If the target's political power relies on the correct operation of systems that are not only electronically isolated but also hidden, penetration into lesser systems may leave little impression on the target. Note that penetrating a system and persisting within it require similar skill sets but different technologies. Penetration requires knowledge of vulnerabilities; persistence requires knowing how to evade intrusion and anomaly detection systems.

Does the ability to break into a system prove the ability to break a system? From a technical perspective, no. Contrary to some assertions, the ability to read files does not imply the ability to write to them, hence to alter them, just as the ability to watch Netflix videos on a laptop does not imply the laptop's ability to edit such videos. Breaking a system requires not

only breaking into administrator (or otherwise privileged) accounts but also figuring out how to induce a system to fail and keep failing despite many features designed to prevent that. But from a psychological perspective, perhaps the ability to break into a system does prove the ability to break a system—especially if the target leadership cannot discern the difference between breaking into and breaking. If the penetration—a violation, as it were—comes as a shock, the prospect of further implications may startle the leadership—regardless of how technically unfounded such implications are.

Breaking a system, however, is a more hostile, and more difficult, act than breaking into a system. It requires understanding the characteristic failure modes of the system. Creating necessary effects also requires shaping the attack so that the target's system administrators cannot detect and repair the system very quickly, with the definition of "very quickly" being necessarily specific to the context. An attack on a logistics system might have to last days or weeks before crippling its user population. An attack on surface-to-air missile systems, however, only has to disable the systems for the few minutes that attack aircraft are overhead. Nevertheless, it is unclear how fast recovery can be: The history of cyberattacks that require urgent fixes is thin, and documentation from victims of such attacks is even thinner. Perhaps cyberattackers (here and elsewhere) have endeavored to estimate adversary responses by simulating attacks on their own systems and testing their own system administrators' ability to recover their functionality. Even if so, the challenge of conveying to *others* that their attacks can *keep* their systems down for long periods of time is not easy. The intended audiences of such a demonstration may be able to determine what flaw allowed such an attack, fix the flaw, and recover some confidence in their own systems. If so, for brandishing to work, such cyberattack capabilities must be demonstrated repeatedly.

Furthermore, the line between brandishing a capability and employing it can become very thin. Supposedly, the purpose of brandishing is to reduce the other side's willingness to challenge the possessor of cyberattack capabilities. But employing a capability tends to have the opposite effect: It increases the other side's desire to challenge the possessor. It is human nature to hit back. In cyberspace, as with other modes of conflict, brandishing can backfire.

One possible way out of this dilemma is to demonstrate the ability to crash another person's system by demonstrating the ability to manipulate it in ways that, if continued or carried out in other contexts, could demonstrably break it. For example, the demonstrated ability to put a blank spot on a radar during normal operations implies the ability to put one there when the radar is tracking a hostile object. The ability to raise the temperature of someone else's chemical process by one degree may imply the ability to raise the temperature enough to cause serious damage, including destruction of the equipment. Inducing a random blank spot or jiggling the temperature may be hostile attacks but not acts of war. Yet they may suffice to suggest that interference with operations or destroying a chemical facility—which may well be acts of war—are within the attacker's capability. The usual caveats apply. Such demonstrations have to be conveyed to leaders, and such demonstrations are self-limiting if they induce corrections within target systems that complicate repetition. For some systems, jiggling one parameter slightly may not imply the ability to do so dangerously if safeguards exist.

Inducing Fear, Uncertainty, and Doubt

Nuclear arms fostered fear, but there was not a great deal of doubt or uncertainty about their applications.² Cyber may be the opposite—incapable of inducing real fear directly, it may be capable of raising the specter of doubt and uncertainty, especially in the minds of those who might wonder if their military systems and hence their military would work when needed. This would cause queasiness if they had to use force of dubious reliability. The target state need not believe that it will lose a war it otherwise would have won were it not for such implanted logic bombs. To echo Mearsheimer's argument on conventional deterrence,³ it suffices if the potential attacker believes that its odds of winning quickly are not good enough because its systems have been compromised.

An uncertainty-and-doubt strategy may work to the U.S. advantage by persuading other states to be very careful in pursuing a network-centric high-technology force to counter U.S. military capabilities. This means it may be dissuasive. A lot depends on how other states react to the idea that hackers have penetrated their military systems and left behind implants, which, when triggered, could generate rogue messages, alter sensor data, create network dropouts, and even make weapons fail.⁴ It is possible to conclude that, if the target state believes that (1) it has been so hacked, (2) has no alternative but the systems and equipment it has, (3) its estimate of war's outcomes are decidedly worse as a result, and (4) it has a choice on whether to go to war, the state's desire to go to war would decrease.

How might such doubt and uncertainty be induced? The most straightforward way is to hack into such systems and then make it obvious that they have indeed been hacked. Claiming responsibility is unnecessary because the point is to emphasize not U.S. power but the vulnerability of targeted systems to cyberattacks in a way that leaves their owners doubting their own systems. But if the point is not to provide proof but to instill uncertainty, making the result obvious beforehand is unnecessary. In fact, it may be unwise if the first demonstration makes the next one harder to accomplish. Thus, proving a system was, is, and will stay hacked may be impossible. However, the hint of an attack leaves no specific trace and hence no specific fix. Even if system owners react to rumors by making general fixes, such as selective disconnection or the installation of anti-malware guards, there will be nothing that suggests which of these general fixes worked.

In some cases, rumor can be more powerful than fact. After all, it takes, on average, twice as long to find nothing in a room as to find something there. Worse, if finding something is conclusive but sweeping and finding nothing is inconclusive, it takes far longer than twice as long to know that one has found nothing than to find something. System owners may be

² Astute readers may see the term, "fear, uncertainty, and doubt," a phrase Gene Amdahl coined after leaving IBM, to describe the effect that IBM people "instill[ed] in the minds of potential customers who might be considering Amdahl products."

³ John J. Mearsheimer, *Conventional Deterrence*, Ithaca, N.Y.: Cornell University Press, 1985.

⁴ Although the psychological effects of a cyberwar attack are speculative, it may well exceed its real effects. For example, if one just counts the number of centrifuges that destroyed themselves, Stuxnet can account for only a few months' delay in Iran's nuclear program. But, to get a bomb, Iran must commit to enriching uranium from 3 percent (U-235) to 90 percent. During the months required to do this, Western militaries may well react with alarm and force. If Iran cannot convince itself that Stuxnet has not been eradicated, it may conceivably fear that its centrifuges may be ordered to break down in those critical months, exposing Iran to retribution without gaining a bomb for its pains—thereby giving it pause when contemplating going ahead.

unable to rest assured that, having found supposedly rogue code will solve the problem because there is no proof that what was found was the rogue code that rumors referred to; such code could be a glitch unrelated to any malevolent actor or could have been placed there by a third party.

A great deal depends on what others are predisposed to believe about U.S. capabilities with technology in general. U.S. cyberwarriors need never reveal the techniques of this or that manipulation but just ensure there are enough hints out there that say they do have the requisite skills. Subjecting that belief to a test could lead to failure and break the spell they may be under. It cannot be overemphasized that *the target of the attack is not the system itself but confidence in that system and any other system an adversary depends on.*

What helps is the ability to convince others that they cannot protect their systems even after painstaking attention to their security. They may have checked everything three times. Yet the cyberwarriors find their way in. The effect is necessarily prospective rather than retrospective; it is rare these days that people are attacked; the attack makes the news; and yet there is no good idea how the attack was carried out or at least what vulnerability was exploited to enable the attack.⁵ Many of the instruments of the attack remain with the target system, nestled in its log files, or even in the malware itself. Even if the targets of the attack (e.g., the Iranians) cannot figure out what was done or how it was done (e.g., Stuxnet), there may be others who can (e.g., the Belarus firm, VirusBlokAda). The number of prominent attacks whose workings, notably penetration and propagation methods, remain a mystery is small, perhaps zero. To be sure, certain attack methods, notably distributed denial of service, contain no prospective, let alone retrospective, mystery as to how they work; they rely primarily on brute force. Furthermore, anyone who follows the news will understand the ubiquity of hacking. It is no great exaggeration to posit that any information of interest to a sophisticated state sitting on a system connected to the Internet has long ago escaped. At this juncture, there are too many vulnerabilities associated with web scripting (e.g., Java) and document-presentation programs to feel very secure.

The vulnerability of less-sophisticated states to the possibility that others are inside their systems is enhanced when the target does not really understand the technology behind its own weapon systems. Although sophisticated states can be counted on to know military hardware better than unsophisticated states do, the difference is usually a matter of degree. Sophisticated militaries get more from their equipment: An F-16 is likely to be more effective in the hands of an American pilot than in the hands of a typical third-world pilot. Advanced militaries also maintain their equipment better. Still, even an inexpertly flown and indifferently maintained F-16 is a war machine.

An information system, though, may have a negative value in the hands of users unsophisticated or indifferent about security. Poorly defended systems may, under pressure, leak information, buckle unexpectedly, or provide bad data to warfighters and other decisionmakers. In cyberwar, a great hacker can be orders of magnitude more efficacious than a merely good one in ways that do not characterize the difference between a great hardware repairman and a merely good hardware repairman. The difficulty that less-advanced countries have in

⁵ When attack code is encrypted, the decryption process may be very slow if even possible. Part of Stuxnet was encrypted but later broken. As of mid-August 2012, Kaspersky, a major security firm, was unable to break the encryption in the Gauss malware and issued a public call for assistance (Jeff Goldman, "Kaspersky Seeks Help Decrypting Gauss Malware Payload," *eSecurity Planet*, August 15, 2012).

generating impressive cyberattack capabilities may be attributed to poorer educational facilities and a less well-educated recruitment base. Yet their lack of access to others' source code or their not having built any of their own and having few among them who have ever built any operational source code helps ensure their military systems are far more vulnerable to cyberattack than comparable systems of sophisticated states. Third-world nations with turnkey systems are also more likely to be using standard configurations and operating procedures, making attacks on them more predictable than attacks on those who understand their systems well enough to tune them to their unique circumstances. Unless such countries are under official U.S. sanction, their systems could very well be maintained by U.S. firms. If cloud computing comes to match the current enthusiasm of its vendors, critical components of domestic control systems may be stored in other countries and be operated by other entities, of which U.S. firms now appear the most likely hosts.

Would Such a Strategy Work with Russia and China?

With Russia, the answer is almost certainly not. First, Russian capabilities at cyberwarfare are very advanced, as befits a state so devoted to *maskirovka* and blessed with a surfeit of world-class mathematicians.⁶ Russians may fear U.S. military capabilities, particularly in electronics, but are unlikely to regard them as particularly puzzling (especially if electronics are not part of the cyberwar package). Second, Russia's military long suit is not the systems integration of complex electronics and networks. It is precisely because they lack confidence in their conventional military that they lean so heavily on their nuclear arsenal. Thus, it is unlikely that their investment strategy would be diverted by the U.S. development of cyberweapons.

With China, the answer is probably not. China has certainly shown enthusiasm for cyberwar. It appears in their doctrine and in the great volume of intrusions people attribute to them. Chinese talents in cyberspace lean more toward quantity, as befits a focus on cyberespionage (and deep pools of well-trained but cheap labor), than toward the sort of quality required to get into hardened military systems. Furthermore, China's military investment strategy is quite different from Russia's. The Chinese have less interest in achieving nuclear parity and more interest in pursuing anti-access strategies that rely on sensors, surveillance, and missiles, which normally require high levels of systems integration, hence networking.⁷ These factors leave some—but only some—scope for a U.S. dissuasion posture based on using cyberwar capabilities against China.

How the Fear of Penetration Might Affect Enemy Operational Behavior

One purpose in demonstrating cyberwar capabilities is to force states to take the potential for system failure and consequential embarrassment into account and curb their enthusiasm for

⁶ *Maskirovka* is a Russian term meaning "disguise, camouflage, concealment."

⁷ See M. Taylor Fravel and Evan S. Medeiros, "China's Search for Assured Retaliation: The Evolution of Chinese Nuclear Strategy and Force Structure," *International Security*, Vol. 35, No. 2, Fall 2010, pp. 48–87, and Roger Cliff, Mark Burles, Michael S. Chase, Derek Eaton, and Kevin L. Pollpeter, *Entering the Dragon's Lair: Chinese Antiaccess Strategies and Their Implications for the United States*, Santa Monica, Calif.: RAND Corporation, MG-524-AF, 2007.

war. But would it? Perhaps not. First, when it comes to war, nearly all defenders and a surprising percentage of attackers believe that they have been put into a position where they have no choice but to wage war because the alternative is worse, e.g., fighting later would put them at a greater disadvantage—so the Japanese believed in 1941 or the Germans in 1914. Fear has already failed to deter them. Second, how badly do countries contemplating such actions need high-technology systems to succeed? Many high-technology systems (e.g., electronic warfare) are needed only against similarly sophisticated opponents, not, say, guerillas. A threat that looks big in peacetime (when systems are vulnerable by dint of being connected) may look smaller in wartime (when systems are configured for survival, in part by being disconnected from the outside world). Finally, the target may simply not believe that U.S. cyber capabilities are good enough to stymie military forces completely—not during peacetime and certainly not when the war drums sound. Going to war requires surmounting a great many fears; the fear of penetration may simply be another.

Persuading third parties that there is a ready-to-go gremlin in their systems carries other risks. At a minimum, if they keep their wits, they will likely pay more attention to operational security after U.S. cyberattack capabilities have been brandished. Any belief that the vector into their systems is a spy will induce them to practice more personnel security. If the winds of alliance shift and the United States has to fight together with such countries, hints of penetration may make it difficult for the United States to work with new partners. Previously benign liaisons with a target country may become more difficult if the “partner” suspects that interacting with U.S. forces reveals how its systems are operated and networked and thus where and how the United States could implant malware in them to the best effect.

Once other states think the United States is behind their fears, reality may be secondary. Countries that are certain that their militaries have been attacked may be less inclined to blame their neighbors, whom they may not credit with enough sophistication to pull off such an attack, and more apt to blame a technologically advanced country, such as the United States or Israel. Indeed, the spread of cyberattack capabilities makes it easy for such countries to hold the United States responsible for *any* failure in military equipment, even for accidents or human error. The instinct to blame others predates cyberspace: Egypt convinced itself, for a few days in June 1967, that the Israelis could not have destroyed its air forces, so the Americans had to have done it. Militaries that can give themselves a pass from their public by using such an excuse may be insulated from the effects of their own mistakes and may maintain their influence and power longer than they should. Alternatively, to the extent that such leaders themselves come to believe their excuses, they may fail to learn from their own mistakes, which may actually help the United States.

Target militaries may also conclude that depending on foreign sources for logic-processing devices is dangerous. This could spur them to build more indigenous production capability or, alternatively, to pressure their suppliers to hand over the source code with the systems. Both can be negatives for the United States to the extent that they are currently being supplied by U.S. corporations. The same suspicions may color the target’s agenda toward civilian gear, such as routers, used in their networks. In response, they may pursue indigenization, more-transparent source code, and better cyberdefense training. If they convince themselves that adherence to the Windows/Intel standard is the root of the U.S. ability to hack their systems, they may lean toward more-open operating systems or make common cause with other

countries, such as China, that are striving to build a foundational layer from components and code not believed to be controlled by U.S. companies.⁸

The problem does not go away if the hints that other systems have been penetrated turn out to be baseless. Assume that the United States has convinced others that it can interfere with anyone's military equipment. Then a war breaks out and no equipment fails in an unexplained way. Observers will conclude that the United States chose *not* to disrupt the sophisticated systems of one side. If only one side's equipment works, others may assume this to be proof that the United States must have played favorites and even blame the United States for atrocities associated with such military equipment. They may not pay attention to counterarguments: The United States hinted "might" not "would"; it cannot afford to get into everyone's equipment; some equipment is inaccessible to the outside; other equipment, such as AK-47s, simply has no electronics to get into. Until the hints started flying, no one could imagine that military equipment could be remotely disabled—but afterward, no one could imagine the United States *not* being able to do it.

How Fears of Penetration Might Affect Defense Investments

A state afraid of penetration could pursue compensatory strategies. It may observe that the effects of cyberattacks are temporary and difficult to repeat. It then maintains its investment strategy after convincing itself that, even if its weapons do not work when first used, it can survive the initial exchange and regain efficacy for later rounds of conflict. Such a perspective would have to overlook the ability of high-technology militaries to conclude successful conventional campaigns over the course of days rather than months or years. That is, there may not be a second round. A sophisticated system owner may be able to find and patch a newly exploited vulnerability within hours or days after it has been discovered, especially with outside help. But can an unsophisticated system owner on the outs with the developed world and countering a sophisticated U.S. cyberattack count on so quick a recovery? The state may also realize that, once a system has become ill, warfighters may not want to bet their lives on it until it has been provably cured, a lengthier process than simply having its symptoms relieved.

If states anticipate that their networked systems may be penetrated, they may forswear network-centric warfare. Why try to face foes with weapons that may well fail spectacularly when used? Conversely, for the United States, if it really can defeat the other side's network-centric military capabilities, why would it want to dissuade them from building and depending on them? One reason might be that the United States cannot be certain it can defeat such capabilities but wants to persuade others it can. Another may be that it may want to dissuade a military buildup because it would lead to a more-aggressive security policy and therefore lead it to start or carry on a conflict when U.S. security would be better served by its not trying to use such capabilities rather than by its subsequent defeat when it did. However, if the United

⁸ Iran is even going so far as to disconnect its Internet from the rest of the world's. From Christopher Rhoads and Farnaz Fassihi, "Iran Vows to Unplug Internet," *Wall Street Journal*, May 28, 2011:

On Friday, new reports emerged in the local press that Iran also intended to roll out its own computer operating system in coming months to replace Microsoft Corp.'s Windows. The development, which could not be independently confirmed, was attributed to Reza Taghipour, Iran's communication minister.

See also "Iran to Unveil National OS Soon," *PressTV*, January 4, 2011.

States believed that the other side would go ahead anyhow, it may be better off keeping quiet about its confidence that it can defeat such capabilities.

The target's counterstrategy may be to rely on lower-tech weapons that are robust against cyberattack because they are never connected to anything. So, if U.S. adversaries forgo networking, is an uncertainty-and-doubt strategy thereby defeated or has it triumphed? Would success in dissuading a potential adversary from a high-technology challenge be in the best U.S. interest? Much depends on the kind of wars the United States is worried about. If the goal is to make it very difficult to use conventional forces to defend against invasion or coercion (rather than fight an insurgency), low-technology forces are no match for the United States. Sacrificing quality may provide others the means to pursue quantity, but, so far, the trade-off for others has not been particularly good; quality usually triumphs.

A more subtle counterstrategy is to network warfighting machines that stay off the Web and forget about networking people. This has the advantages of permitting air-gapping as a defense strategy and avoids some of the vulnerabilities arising from human error (notably those associated with authentication, such as passwords and tokens). If networking warfighters is oversold, states that forgo it may be doing themselves a favor. Or they can network their equipment together but snip their links to the rest of the world. Perhaps a self-denial-of-service policy reduces their military's ability to learn from others and, to some extent, itself. Yet, many militaries are so self-contained that, even in the absence of cyberwar, they are apt to discount the experience of others from whom they might learn something.

The Algebra of Direct Intimidation

If announcing offensive capabilities fails to deter or dissuade, might a demonstration be worthwhile to create a coercive capability?⁹ One dilemma lies in how far to take credit for any demonstration. Consider the following algebra. Assume that, if an attacker, call it state Z, reveals itself unambiguously through its cyberattack, it loses more from retaliation than it gains in coercion. If it hides itself absolutely, it gets no benefit from coercion (the damage might easily have been an accident). It would seem that intermediate levels of assurance yield linear net negative benefits. For instance, if the target thinks that the odds that the attacker was state Z are 50:50, the coercive benefits are half of what they would have been were the target certain.¹⁰ Similarly, the odds of retaliation—and thus the expected cost of bearing such retaliation—are

⁹ One method of demonstrating cyberattack capabilities is to attack a state that clearly deserves it and use its fate as a lesson for others. Such a state should be one that relies on some infrastructure and is not very good at protecting it. It helps if the target state is generally not sympathetic and has no good option for responding without escalating matters more than it is prepared to handle. Overall, however, there are more than enough reasons to recommend against trying this. The effect requires some attribution, at least implicit, on the attacking state's part—otherwise the only thing being demonstrated is that some states build infrastructures they cannot defend. But such a policy makes the attacker look like a bully. It also legitimizes cyberwarfare. Other states may be impressed by the attacker's chutzpah but not necessarily its acumen. It is too easy for those one would impress to counter that they are hardly as vulnerable as the state that was attacked. If the attack is permitted by a weakness that others shared, they may take the results of the attack more seriously—but only long enough to fix similar vulnerabilities of their own.

¹⁰ If the target state thinks that the odds that state Z would carry out a second, perhaps more consequential, cyberattack in response to something it might do are 50:50 (that is, precisely equal to the odds that it thinks state Z carried out the cyberattack), it would weigh the expected cost to itself of a reaction from state Z should it go ahead and do so half as heavily as it would have if it were certain state Z carried out the cyberattack.

half of what they would be were the target certain. Thus, from state Z's perspective, both the benefits of coercion and the expected cost of retaliation are halved. This still leaves a net negative. So, it appears that it cannot win.

But are the odds of retaliation really the same as the perceived likelihood that Z was the attacker? In more-specific terms, are the odds of retaliation a 50:50 proposition if the target thinks the odds are only 50:50 that Z was the attacker? A great deal depends on how risk-averse the target is. If the target fears the consequences of not retaliating against the true attacker (the wimp factor) more than it fears the consequences of retaliating against an innocent state (the blunder factor), it does not need much confidence in its attribution to convince itself to hit back. What seems more likely is that the target fears the blunder factor more than the wimp factor, and a 50:50 confidence level will not be enough to persuade it to retaliate. In that case, the odds that the target will retaliate when it is only half sure that state Z did it would be less than 50:50.¹¹

If so, the coercive force of a cyberattack when the target thinks state Z might have done it—and thus the benefit of coercion to state Z—may be greater than the expected cost of retaliation. The benefits of coercion scale with the degree of confidence the target has that state Z did it. Yet the cost of retaliation only has to be taken into account when attribution is sufficiently clear that the target thinks the odds of making a mistake are sufficiently low.

The broader lesson is that an attack that might be but also might not be attributable may be worthwhile for the attacker, even when a more obvious attack is not. The target state, for its part, may do its best to exaggerate its likelihood of retaliating, the better to throw off the attacker's calculations. Yet, given the nature of crises and the natural ambiguities of cyberspace, the attacker is likely already dealing with much ambiguous information.

If the attacker's coercive objectives are more general and do not depend so much on who the target thinks the attacker is, its net gain is even larger. Telling another "do what I want" without identifying "I" is hard—but not impossible. Suppose a country (e.g., an Islamic state) has allied its interests with a larger community's (e.g., the *umma*'s), particularly one with powerful nonstate actors. If so, some correlation can be made between the timing and nature of the attack (e.g., following action against Islamic individuals) and the behavior required of the target (e.g., stop attacking Islam) without necessarily indicting the attacking state. A state accused of a cyberattack could plead that it has friends that it cannot control but whose righteous ire should be acknowledged. So-called patriotic hackers may be citizens of an accused state without that state appearing hypocritical as long as it makes a credible attempt to bring them under ostensible control. Alternatively, the state can take satisfaction in cyberattacks that punish behavior that contravenes the community's interests. At the same time, it need not admit to the support, much less to the protection or even sponsorship, of such attacks. The attack's coercive potential may be limited to the values held by the community—normally just one among its overall interests (e.g., what may be good for taking action against a common enemy may not be so good for asserting particular interests, such as water rights). But that may be enough.

¹¹ To illustrate as much, assume the target thinks that the attacker is as likely to be state Y as it is state Z (but does not believe that Y and Z colluded). If it retaliates against one, why not against the other? The only way that could be justified is if the target believes the consequences of hitting an innocent state Y are worse than those of letting state Z get away with an attack.

Would the behavior of the target state ever evolve in the direction the attacker desires as a result of coercion (a question relevant to kinetic threats)? Assume two things. First, attacks that yield less pain than some sensitivity threshold are too weak to coerce the target state. Second, attacks that yield more pain than some response threshold induce the target state to hit back or at least turn less cooperative (at least overall, if not necessarily on the point at issue). If the sensitivity threshold is less than the response threshold, there may well be a zone in between where the target yields to the attacker's will. But if the two are reversed, no attack, however carefully calibrated, will change the target's policy in the desired direction. The attack will be either too weak to be sufficiently coercive or too strong to be absorbed without response—and maybe both. The United States has dramatically demonstrated at least twice that it reacts harshly to being attacked, in response to both Pearl Harbor and the September 11, 2001, hijackings.¹² Granted, the first may not have been an act of coercion (since Japan believed it was going to have to fight the United States sooner or later anyway), and the second may have been carried out to goad the United States into intervening in Afghanistan. Yet, such distinctions aside, the United States proved that coercing it may not be particularly useful when the target's response threshold is lower than its sensitivity threshold. A large body of literature on coercion shows how difficult it is to compel states to comply with demands, even with kinetic weapons.¹³ It is hard to argue that cyberweapons, with all their uncertainties, would do a better job.

The attacker could carry out a covert coercion campaign using *sub rosa* attacks. That is, it can go after systems whose failure or corruption may be costly to the target government but whose effects are not obvious to the public. By doing so, the attacker gambles that the positions of policymakers' sensitivity and response thresholds differ from those of the public. Policymakers, feeling pain and unforced by public opinion, may be freer to yield to coercion, especially if yielding is also invisible to the public.

Direct intimidation may work better if a cyberattack is clearly structured to damage much less than it could have.¹⁴ All attempts at coercion evoke in its victim a mix of anger for having been hit and fear of the next hit. If the initial attack is mild, the anger component may be assuaged by the fact that, while the insult is clear, the injury is not. The fear component, however, is just as great with a pulled punch as with a fully formed punch—as long as the target understands that the punch was, in fact, pulled (although in the ambiguities of cyberspace, the clarity of that message could be lost).

¹² The United States even reacts harshly when it thinks it has been attacked, even if later facts suggest otherwise. The Spanish found this out after the USS *Maine* was sunk in 1898—by what is now believed to have been an accident and not a mine. That noted, a harsh response is not a guarantee, as the lack of response to the 1968 capture of the USS *Pueblo* and the Iraqi missile attack on the USS *Stark* showed.

¹³ Among those who have made similar arguments are Robert Pape (in Robert A. Pape, *Bombing to Win: Air Power and Coercion in War*, Ithaca, N.Y.: Cornell University Press, 1996); David Johnson (in David E. Johnson, Karl P. Mueller, and William H. Taft, *Conventional Coercion Across the Spectrum of Operations: The Utility of U.S. Military Forces in the Emerging Security Environment*, Santa Monica, Calif.: RAND Corporation, MR-1494-A, 2003); Karl Mueller (in Karl P. Mueller, Jasen J. Castillo, Forrest E. Morgan, Negeen Pegahi, and Brian Rosen, *Striking First: Preemptive and Preventive Attack in U.S. National Security Policy*, Santa Monica, Calif.: RAND Corporation, MG-403-AF, 2006), Daniel Byman (in Daniel Byman, Matthew Waxman, and Eric V. Larson, *Air Power as a Coercive Instrument*, Santa Monica, Calif.: RAND Corporation, MR-1061-AF, 1999); and Forrest Morgan (in Forrest E. Morgan, Karl P. Mueller, Evan S. Medeiros, Kevin L. Pollpeter, and Roger Cliff, *Dangerous Thresholds: Managing Escalation in the 21st Century*, Santa Monica, Calif.: RAND Corporation, MG-614-AF, 2008).

¹⁴ See, for instance, Thomas C. Schelling, *Arms and Influence*, New Haven, Conn.: Yale University Press, 1966, notably Chapter Three.

Paradoxes of Intimidation

Are the short-term gains from this sort of intimidation, even if latent, worth the long-term discomfort from accelerating the evolution of a particular class of weaponry? In the nuclear race between the United States and the Soviet Union, Khrushchev would boast that his country could turn out missiles “like sausages.” The United States reacted by accelerating its own missile program. By 1961—a year before the Cuban Missile Crisis—the United States knew it had a strategic edge in nuclear delivery systems, notably missiles. Similar Soviet perceptions persuaded them to ship missiles to Cuba to adjust the strategic balance. The Soviet reaction to having to back down in Cuba was to accelerate its own programs to achieve parity, which they did, thus setting the stage for the Strategic Arms Limitation Talks. Perhaps, had neither side flaunted its capabilities, the same parity and negotiations might have arrived at roughly the same time but at much lower levels. The missile race is hardly unique, as the pre–World War I Anglo-German shipbuilding rivalry demonstrated.

Nevertheless, a cyber arms race is not the most likely course of events. In great contrast to most military weapons, the damage from a cyberattack tends to reflect the characteristics of the target more than the characteristics of the weapon. So the competition to reduce vulnerabilities may overshadow the competition to find and exploit them. Even were this not so, either side’s cyberweapons’ capabilities are a matter of serious dispute—an observation that undergirds this whole discussion. The sorts of numbers that inform the balance of missiles or dreadnoughts (World War I-era battleships) have no meaningful equivalents in cyberspace.

U.S. Policy and the Legitimization of Cyberwar

The current U.S. posture on cyberweapons is coy; it stands between U.S. posture on nuclear weapons (terrible—but useful for very special occasions) and chemical and biological weapons (too sinful even to contemplate). Although there is no official policy affirming that the United States *would* use cyberattacks, there has not been much refutation of the proposition that the United States was responsible for the Stuxnet attacks.¹⁵ Similarly, the press reported, again without refutation, that the Defense Advanced Research Projects Agency (DARPA) was work-

¹⁵ The best example of such an argument is David Sanger, “Obama Order Sped Up Wave of Cyberattacks Against Iran,” *New York Times*, June 1, 2012, p. 1.

In an earlier interview, Melissa Lee asked then–Deputy Secretary of Defense William Lynn about this. According to one report (Kim Zetter, “Senior Defense Official Caught Hedging on U.S. Involvement in Stuxnet,” *WIRED*, May 26, 2011), Lee asked Lynn outright:

“Was the U.S. involved in any way in the development of Stuxnet?”

Lynn’s response is long enough that an inattentive viewer might not notice that it doesn’t answer the question.

“The challenges of Stuxnet, as I said, what it shows you is the difficulty of any, any attribution and it’s something that we’re still looking at, it’s hard to get into any kind of comment on that until we’ve finished our examination,” Lynn replies.

“But sir, I’m not asking you if you think another country was involved,” Lee presses. “I’m asking you if the U.S. was involved. If the Department of Defense was involved.”

“And this is not something that we’re going to be able to answer at this point,” Lynn finally says.

ing on developing offensive cyberwar capabilities.¹⁶ The United Kingdom and Canada have announced similar sentiments.¹⁷

Other nations, such as China, do not even go that far in terms of admitting their cyber-attack capability and why they might use it; yet, their blanket denials are not seen as credible. This only partly reflects the great range of opinion on whether cyberattack capabilities are weapons of mass destruction or even weapons of mass disruption. The attacks on Estonia proved only that they can be weapons of mass annoyance. On a day-to-day basis, cyberattacks are treated as criminal matters, which effectively delegitimize them as tools of statecraft. It is less clear whether such delegitimization is associated with the act per se or its use by private parties. The United States would prefer that other nations treat cybercrimes with greater seriousness and avoid the temptation to privatize the application of military force in cyberspace (e.g., linkages between the Russian government and the Russian Business Network or between China and its freelance and “patriotic” hackers). Ironically, moves to legitimize such weapons may make it easier for other countries to take ownership, hence responsibility for their peoples’ use of such weapons.

All that noted, cyberwar has probably already passed the legitimization threshold. It may have done so, at least against military targets, back in 1999.¹⁸ The United States and similarly capable countries are discussing efforts to delegitimize the use of cyberwar against certain classes of targets (e.g., hospitals). International consensus or even a treaty may result. If so, brandishing a capability to cross these norms would be problematic.

¹⁶ Ellen Nakashima, “With Plan X, Pentagon Seeks to Spread U.S. Military Might to Cyberspace,” *Washington Post*, May 30, 2012, p. A1. The DARPA announcement of an industry day associated with this program, however, noted: “The Plan X program is explicitly **not** funding ... cyberweapons generation” (DARPA, “Cyber Experts Engage on DARPA’s Plan X,” press release, October 17, 2012).

¹⁷ According to “Cyber Strikes a ‘Civilized’ Option: Britain,” *Technology Inquirer*, Agence France-Presse, June 3, 2012:

Preemptive cyber strikes against perceived national security threats are a “civilized option” to neutralize potential attacks, Britain’s armed forces minister said Sunday. Nick Harvey made the comment at the Shangri-La Dialogue security summit in Singapore in relation to reports that the United States had launched cyber attacks to cripple Iran’s nuclear program. . . . Britain’s stance was supported by Canadian Defence Minister Peter Gordon MacKay, who likened a pre-emptive cyber strike to an “insurance policy,” warning of the need to be prepared.

¹⁸ Department of Defense, Office of General Counsel, *An Assessment of Legal Issues in Information Operations*, May 1999.

Brandishing Cyberattack in a Nuclear Confrontation

It is not easy to confront countries that threaten to use their nuclear capabilities if the United States does not conform to their wishes. Might brandishing a cyberattack capability influence the course of such confrontations? Examining the mechanisms of such influence may shed further light on the opportunities and limitations of brandishing cyberattack capabilities (even if not necessarily expanding our understanding of nuclear confrontations as such).

In doing so, we will not necessarily claim that U.S. cyberattack capabilities can reliably confound adversary nuclear capabilities. States, after all, pay a great deal of attention to their nuclear weapon systems against the day when their regime's survival rests on the weapons' being ready for use. Nuclear states go to extraordinary lengths to protect their command and control over such weapons. Nuclear weapons and the large rockets that carry them largely preceded the digital revolution, and the weapons remain largely analog despite the later development of command and control and accurate missile guidance, which do have digital elements.

However, the possibility that the United States *could* penetrate the command, control, or operations of nuclear systems cannot be easily disproven either. Iran's leaders undoubtedly thought that the isolation of their Natanz centrifuge facility rendered it safe from cyberattack. Then they learned about Stuxnet.

Our inquiry is therefore more humble. Could a U.S. threat that it *might* interfere with a rogue state's nuclear weapon delivery help shape a nuclear confrontation? For this question, assume a rogue nuclear power with a handful of weapons capable of hitting nearby countries (but generally incapable of hitting the continental United States). The United States has a robust cyberattack capability (in general terms), from which the rogue state's nuclear arsenal is not *provably* immune. Although the United States enjoys escalation dominance, the rogue state is far more willing to go to the nuclear brink than the United States is. The rogue state (thinks it) has more at stake (i.e., regime survival). Furthermore, it may act in ways that are irrational by Western perspectives.

We first model a two-state confrontation, then later introduce a friendly state on whose behalf the United States has intervened. The United States enters this scenario facing the choice of acting when doing so risks the rogue state releasing a nuclear weapon. Whether the threat is explicit or implicit is secondary. The usual calculus applies. The rogue state is better off if its threat leads the United States to stop. The United States is better off ignoring the threat and going ahead with what it would have done in the absence of the threat *if* the threat can be

nullified but cannot *know* that it will be for certain. The rogue state understands that if it *does* use nuclear weapons, it could face great retaliation.¹

If the United States acts (successfully) in the face of warning and if the rogue state does not use nuclear weapons, the United States achieves its objectives *and* wins the overall confrontation.² If the United States flinches, the rogue state wins. If the rogue state uses its nuclear weapons and if, as is likely, the United States responds likewise, the rogue state loses greatly, but the United States is also far worse off.³

Two-Party Confrontations

In a confrontation in which disaster would result from both sides carrying out their threats, each must ask: Are such threats credible? If one side thinks the other will yield, it pays to stand firm. If it thinks, however, that the other is implacable, it may have no good choice but to yield itself. The projection of implacability is beneficial, but the reality of implacability is frequently suicidal.

Note that the basis for the implacability can also be entirely subjective, which is to say, unfounded on the facts of the matter. If one party is convinced that it will never pay a high price for being implacable, communicates as much, and acts as if it were so, the other cannot take any comfort from the fact that the first has no technical basis for the belief. The only consideration is whether the first party actually believes as much, is willing to act accordingly, and can ignore the logic that whispers that no one can possibly be completely confident on the basis of iffy information. To one party, the willingness to act on the basis of the impossible seems like cheating. To use an analogy, imagine a game of “chicken” in which the driver of one of the two oncoming cars throws the steering wheel out the window. This cheat forces the opponent to choose between a certain crash or veering away (and thus losing). However, when the consequences of a crash are far greater than the benefits of winning, this strategy is irrational if there is a nontrivial likelihood that the other side will be intent on punishing cheaters at the cost of

¹ If the United States were given a choice between having a nuclear deterrence capability and being able to brandish cyberattack capabilities, the first would clearly be preferable. Nuclear deterrence already exists, yet the U.S. strategic community worries about rogue states having nuclear weapons, and U.S. policy goes to great lengths to prevent Iran from acquiring its own weapons. This suggests that there is less than complete confidence that the United States could always deter nuclear use under all circumstances. It is in that context in which we ask: Might brandishing cyberattack capabilities be of some assistance here?

² The focus on “winning” a nuclear confrontation is not an argument that the United States should use such tactics as the sole or even primary method of defusing the threat from nuclear rogue states. Much can be said for precrisis policies to remove the incentive or ability of rogue states to acquire nuclear weapons, for crisis policies that attempt to persuade the rogue state that it ought to follow international norms of conduct, and for policies that give the rogue state an honorable exit even from a dilemma of its own making.

³ For the sake of tractability, this analysis will ignore many of the options and branch points that exist in even the simplest nuclear confrontation of this sort. For instance, if the threat from the nuclear rogue is implicit, it may not be obvious that any U.S. act crosses the line. Under such circumstances, the loss of face that the rogue state would experience by not responding is less. However, if the United States concludes the implicit threats ring hollow once, it may feel it is home free. If the rogue state wants to preserve its ability to threaten, it may have to find a second threshold (or what is more difficult, try to compel the United States to pull back or avoid repetition) and shift to an explicit threat. Similarly, there may be multiple thresholds of nuclear use, some of which may invite all-out retaliation and others not. Options may include, in order of severity, a demonstration shot, a burst that destroys equipment but not people (e.g., an electromagnetic pulse), an antiship or antifleet attack, an attack on ground forces, an attack on an allied population center, and an attack on U.S. soil.

all other values. In the analogy, the second driver might rather crash than lose to a cheater.⁴ But in general, a strategy of implacability, can, if credible, do well, as long as the other side is not equally implacable.

So, the United States creates the belief (whether by saying so, hinting, or letting others draw their own conclusion) that the rogue state cannot carry out its nuclear threat. That is, the United States acts as though a flaw somewhere in the nuclear command-and-control cycle, probably an induced flaw, prevents immediate nuclear use. A lesser case is that the command and control is less certain, the weapon is weaker, and/or the delivery system is far less accurate than feared.⁵ Although permanently disabling a nuclear command-and-control system is quite a stretch for cyberwar, it is less fantastic to imagine that the United States could delay a weapon's use. A temporary advantage, though, may still give the United States time to cross the red line and thereby attain a *fait accompli*.

So posturing, the United States prepares to cross the red line, while communicating its confidence that the rogue state will not retaliate. This confidence stems from a combination of its own nuclear deterrence capability *plus* its ability to confound the rogue state's nuclear capability: The rogue nuclear state probably will not decide to retaliate, and if it did decide to, probably *cannot* retaliate. The *combination*, in this case, is what reduces the odds of a nuclear response to a sufficiently low level, if the rogue state is at all rational. Even if it later assures itself and others that its nuclear capacity is intact, but the United States has already acted, the onus then falls on the rogue nuclear state to respond to what could well be a done deal. If the rogue state understands the logic *before* brandishing its own nuclear weapons, it may choose not to ratchet up tensions in advance of the U.S. crossing red lines.

This strategy requires the rogue state to *believe* that the United States is implacable—based, in part, on the possibility that the United States *believes* it can use cyberoperations to nullify the nuclear threat. It also helps if the rogue state is not completely sure that this confidence is misplaced, although it may work even if the rogue state believes there is no basis for such confidence, as long as it believes the United States cannot be convinced otherwise.

Note that this implied or expressed belief in a known flaw is somewhat broader than the claim that the United States *caused* the flaw. In many respects, it suffices that the United States knows about the flaw and that the rogue state cannot find it or can do nothing about it over the immediate course of the crisis. The flaw in question could have been created by noncyber means (e.g., a saboteur) or a third party; it could have been there in the design all along. However, a flaw that prevents a nuclear shot is different from a flaw that the United States can exploit to prevent a nuclear shot. The former claim loses credibility if the rogue state can get a shot off, but the latter claim can survive a shot, albeit in weakened form, as argued below.

⁴ But what if one side's determination to make the other side play fair overwhelms the rule of optimizing his own outcomes? This possibility should not be dismissed lightly. People overpunish cheaters even at the expense of their own well-being in gamelike situations. Economists have repeatedly shown as much, notably by watching people play the ultimatum game, in which two players interact to decide how to divide a sum of money that is given to them (Martin A. Nowak, Karen M. Page, and Karl Sigmund, "Fairness Versus Reason in the Ultimatum Game," *Science*, Vol. 289, No. 5485, September 2000, pp. 1773–1775). Indeed, this tendency may be hard wired (see, for instance, Marco F. H. Schmidt and Jessica A. Sommerville, "Fairness Expectations and Altruistic Sharing in 15-Month-Old Human Infants," *PLOS ONE*, Vol. 6, No. 10, October 7, 2011).

⁵ *Nuclear-command-and control cycle*, here, is used loosely to refer not only to the linkage between the order to launch a nuclear weapon and/or detonate a nuclear device but also to the integrity of instructions in the relevant devices themselves. Failures in the latter, for instance, could lead to a misfire, poor aim, the failure to detonate, or premature detonation.

Although the confidence that the United States can frustrate a rogue state's nuclear capability through cyberwar is like being able to do so through, say, effective missile defense, there are important differences. The efficacy of a missile defense system can be demonstrated against an actual missile equivalent in sophistication (e.g., having penetration aids) to the rogue state's missiles. No such tests are decisive in cyberwar because cyberwar generally depends on the target system having flaws that its owner is unaware of. Tests can only be conducted and capabilities demonstrated against systems with particular flaws. Once the nature of these flaws is understood, such flaws can be fixed or routed around. *Afterward, the likelihood that such tests would work against that particular flaw drops sharply.* Furthermore, the existence and some of the basic characteristics of a U.S. missile defense system would be public knowledge and can thus be fed into the public debate over the odds of, say, a successful exoatmospheric engagement and hence the wisdom of facing down a nuclear rogue. Conversely, the existence and the basic qualities of a U.S. cyberattack capability can only be surmised—in large part because offensive cyberwar capabilities must be highly classified to remain effective.

Signaling that the rogue state has a flaw in its nuclear command-and-control cyber does not set up a use-it-or-lose-it dilemma for the rogue state—the vulnerability that enabled the flaw to be introduced necessarily preceded the crisis and, in all likelihood, so did the discovery of the flaw and its exploitation. Thus, the rogue state has no ability and therefore no incentive to use nuclear weapons faster lest it lose it because the cyberattack threat connotes that the rogue state already lost it (but may recover it later).⁶

An even better threat is for the United States to suggest that it will respond to a failed nuclear launch just as it would to a successful nuclear launch. This would inform the rogue state that it may lose a great deal by attempting a nuclear shot. This might convince the rogue state to back down—if the rogue state has any basis for believing that such a shot might not work. Failure would lead to reprisals as devastating as success would have. Worse, the credibility of the rogue state's capability would have been sharply reduced. But can the United States retaliate for a failed launch? It is one thing for a nuclear shot to fail visibly; everyone else can read intent. But if a missile fails to launch, can the United States retaliate, especially with nuclear weapons, on the basis of evidence that can only come from cyberspace and thus may be perceived as concocted?

In the real world, while nuclear weapons tend to be analog and delivery systems a combination of analog and digital, the permissive action links that prevent accidental launch of nuclear systems (particularly those of the newer nuclear states) are digital. If newer nuclear

⁶ A RAND colleague surmised: Why bother brandishing a threat that the United States can render a rogue's nuclear command and control inoperable instead of actually rendering it inoperable? That is prevention, not brandishing. If the answer is that the United States might not be able to prevent a launch but wants the rogue to believe that it is able to do so, the United States must bluff. Because it is clearly better, given the nuclear stakes, to prevent than merely to display one's ability to prevent, it follows that brandishing implies bluffing.

Given this logic, a rogue would assume that any U.S. cyberattack that did not in fact prevent a nuclear launch was a bluff. An initial response to that line of logic is that, by brandishing the capability to stop the rogue state's nuclear cycle, the United States is declaring that it has *already* prevented launch and just wanted the rogue state to know as much so that the rogue state, hitherto confident that it had the ability to carry out a nuclear launch, did not, as a result, put itself in a position where it could not back down and might even have to double down. The rogue state may then reason that U.S. brandishing of such a capability could come only at a cost to the United States because it would raise the likelihood that the vulnerability that permitted the cyberattack could be found. To take the brandishing seriously (rather than as a bluff), the rogue state would have to calculate that the United States believes defusing the crisis before the rogue state works itself into a position that could be perilous for both sides outweighs the possibility of the disclosure imperiling the cyberattack.

powers suspect that the United States is trying to thwart weapon use by interfering in the digital components of nuclear command-and-control systems, they may conclude that the United States had found some way to interfere with permissive action links. It would not serve stability for countries to start disabling them out of fear that they may have been tampered with. If countries do disable these links, the odds of an accidental launch rise. Hence, in the real world, the threat to use cyberattacks against nuclear operations of rogue states has important and negative ramifications in the longer run.

Disabling a Capability Versus Thwarting a Threat

Thwarting the rogue state's ability to *threaten* nuclear use is different from thwarting the rogue state's ability to *use* a nuclear weapon. Thwarting a threat requires projecting confidence to the adversary that the capability behind the threat will not work—risking the possibility that the rogue state, so alerted, will reexamine its nuclear command-and-control systems and either fix the flaw or route around it (e.g., allow more opportunities for “manual” overrides of an errant electronic system). Thus, today's hints might reduce the likelihood of actual compromise over time.

Conversely, having a cyberattack capability and *not* brandishing it will not relieve the pressure on the United States to withdraw from a crisis in which a rogue state brandishes, but with nuclear weapons. Such pressure may come from internal opposition; allies; respected interlocutors; or, as discussed below, the very country for which the United States has intervened to defend against the rogue state. Even then, having the capability is not useless, even in the public arena. If the knowledge that the foe's nuclear capabilities can be neutralized gives U.S. leadership enough confidence, others may infer that the United States has valid reasons for its confidence, even if these are never revealed. For some, that would be good enough.

The tension between winning a confrontation by conveying the adversary's weakness and limiting war damage by exploiting the adversary's weakness suggests the need to weigh the odds that the rogue state would use its weapons. If the odds of use are low (e.g., because U.S. implacability is expected to have the desired effect), more hints are warranted; if the odds of use are high, more silence is the better option. If the rogue state understands as much, it will conclude that, if the United States *is* trying to convey that it knows a secret, the United States believes the likelihood of actual use is low because it can, in fact, deter the rogue state. This has the effect of reinforcing the message that the rogue state's nuclear flaws should be taken seriously or at least that the United States is taking them seriously.

The Rogue State Might Try to Discredit the Cyberwar Bluff

Assume the rogue state, having scrubbed its nuclear command-and-control systems, manages to brush aside all its doubts about their reliability—despite the difficulty of being certain that the (purportedly all-knowing, all-seeing) U.S. intelligence community does not know something about the state's nuclear systems that it has overlooked (remember Stuxnet) and being mindful that surprises in cyberspace are surprising in their detail, yet seem to occur frequently.

If the portrayed basis for U.S. confidence were contradicted by what leaders knew about their nuclear command-and-control system, it may be very hard to convince the rogue state that the United States felt truly confident nonetheless. This logic may motivate the nuclear rogue state to convince the United States that its ability to carry out the nuclear threat would not be impaired by a cyberattack.

But how? One way is to reveal something about its command and control that would indicate that the basis for the U.S. belief is illusory. Yet, the more it reveals about its command and control, the more information it gives for U.S. (and every other country's) cyberwarriors to work with in pursuit of a new flaw. Such revelation may also give useful information to adversary special forces, as well as targeting information to adversary air forces. There may also be some hesitation if the rogue state feels that revealing secrets about command and control may tell *internal* audiences (e.g., some members of its military) things that its leadership prefer they not know.

The other way is by a successful launch and detonation. Success would eliminate the possibility that the induced or discovered nuclear command-and-control flaw is pervasive and endogenous. To be sure, a failed shot signals a willingness to use such weapons in ways that mere possession does not. A successful shot, however, does not eliminate the possibility that the exploitation of the flaw was one the United States can activate whenever it chooses or one that appears only under certain launch parameters (e.g., if the aimpoint of the missile is in territory the United States does not want hit). The United States, after all, could have refused to interfere with the nuclear strike in the belief that the results had no grave consequences (e.g., it was a demonstration shot in the rogue state's territory). This still leaves open the possibility that, if the shot *did* matter, the United States could have stopped it. If so, the United States could still maintain its determination to cross whatever red line the rogue state established.

The rogue state, anticipating as much, would have to make its point with a nuclear shot that would generate consequences (1) too weak to merit a devastating response from the United States but (2) strong enough to exceed what the United States could tolerate with equanimity (note the similarity to the presumed anger-fear interval discussed above). It would carry lower risks for the nuclear rogue state (because retaliation would not follow) but would discredit the basis for U.S. confidence (because if it could have stopped the shot, it would have done so). One possible example is a nonnuclear warhead atop a nuclear delivery system shot toward a destination where detonation, if nuclear, could have killed many. In theory, the United States could not have known that the shot was nonnuclear until detonation. In retort, the United States could point out (more likely hint) that the level of knowledge required to get inside the nuclear command-and-control cycle is also more than enough to distinguish a nuclear from a conventional shot. As a more practical matter, if the United States does not establish the line at which it would block a shot and the line at which it would retaliate devastatingly, how would the rogue state know if there was any daylight between the two, much less where it was?

Conversely, a conventional demonstration shot (1) gives others reason to question the need for a demonstration (perhaps confirming that rumors of a flaw have a basis), (2) uses up at least one missile from a small stockpile, and (3) may not address the suspicion that a flaw in the nuclear weapon's fusing device (physics package) is still in play.

The U.S. strategy also needs careful consideration if the rogue state can start to build a case that its nuclear systems work. The rogue state may not have a constituency to answer to, but the United States does, particularly if U.S. lives are at stake and, even more, if its actions involve allies (discussed later). Even if the United States expresses confidence in its words and actions, how does it convey the basis for its confidence to others? What does it reveal to do so? In pondering this question, strategists have to know that the rogue state may contradict any specific revelation and, even if the revelation were valid, could address it in time if the claim is specific enough to determine where to look for the flaw.

Would it be useful to persuade the rogue state that the United States *could* have such control over and above convincing the rogue state that the United States *believes* it does have such control? It helps if the rogue state's leaders are convinced that the United States *could* have such control *if* they know they do not understand the software built into their nuclear command and control well enough to know that it was invulnerable. One advantage to persuasion is that, if the rogue state's leaders think they would be firing blanks, they would see themselves as trying to bluff and would therefore approach the crisis with a greater tentativeness. If they fear the United States would detect this tentativeness, they might reason that the United States would approach the crisis *as if* the rogue state were bluffing, which, again, puts the rogue state in the position of losing if it backs down but losing much more badly if it launches a nuclear attack.

Ironically, the more the United States wishes to convey its confidence, the less credible such confidence would be. If the United States were truly confident that it could stop a nuclear attack, it would not care very much that the rogue state felt otherwise; either way, nothing bad happens. Thus, wanting too badly to project confidence implies the United States cares what the rogue state thinks—which means that its confidence is less than complete.

Can Cyberattack Brandishing Forestall Unilateral Nuclear Use or Threat of Use?

Can brandishing cyberattack capabilities be useful when U.S. steadfastness is *not* a relevant issue? Examples may be precluding a unilateral nuclear shot (e.g., a bolt from the blue) or persuading a rogue state not to raise the nuclear ante or work itself into a position where it felt it had no choice but to use nuclear weapons. In the latter case, a determined U.S. stand, even one backed by brandishing cyberattack capabilities, could end poorly. Is there a point to signaling to a rogue state that its nuclear command and control may have flaws when the rogue state is not under great urgency to use that capability (as it would be if it tried to stop the United States from crossing a red line)?

Unfortunately, the earlier a rogue state is alerted that its nuclear command and control may have flaws, the more time it has to find and fix them.⁷ Whether the rogue state can *then* convince itself, *after the fix*, that its command system has actually been fixed is a different issue. The number of problems in its system is clearly one fewer. If it thought it had only one problem before, it has zero now—so back, with confidence, to the brink. On the one hand, how would it know that it started out with just *one* exploitable flaw? The act of finding one flaw may just be an indication that it had multiple flaws to begin with. Thus, the *expected* number of exploitable flaws could well be greater after having eradicated one than it was beforehand.⁸ Conversely, if

⁷ True, a cyberattack that can break something would have a longer effect, but it is much easier to interfere with a complex command sequence than it is to insert a destructive command into a complex command sequence with fail-safe features. That Stuxnet destroyed centrifuges without seriously denting Iran's uranium production suggests that these centrifuges were in their start-up phase. Machinery in this phase tends to be reprogrammed more frequently and, hence, is more subject to induced programming errors than machinery that has been running consistently for a long time.

⁸ This is perfectly consistent with the Bayesian model of inference. Assume that the rogue state believes that the number of exploitable flaws in its system is a variable with an a priori probability distribution as follows: an 80-percent likelihood that there are zero flaws, a 10-percent likelihood that there is one flaw, and a 10-percent likelihood that there are two flaws. The expected number of flaws is 0.3 (80 percent times 0 plus 10 percent times 1 plus 10 percent times 2). A flaw is then

the rogue state searches and finds no flaw, it may conclude either that the United States was bluffing or that U.S. cyberwarriors were so clever and subtle that the relevant exploit was undetectable even after a painstaking search.

Friendly Third Parties Add Complications

Friendly third parties that can veto U.S. actions can complicate brandishing cyberattack capabilities to bolster a U.S. stand. These parties have many ways to exercise their veto, not least by denying U.S. forces access to their territories. Even if the U.S. military can operate from the sea or distant air bases, third-party objections to pending U.S. actions could undercut the primary U.S. rationale for any action in the region.

Although friendly third parties may well have greater animus against the nuclear-armed state and a correspondingly greater willingness to see it humiliated, perhaps disarmed, they may well blanch at the cyberwar bluff. First, they and their citizens are likely to be at greater risk, by dint of sitting within range of the rogue state's nuclear systems. Second, they would know less than the United States does about what cyberwar capabilities the United States has and how they might be used. As a result, third parties may have less confidence that such plans would work.

None of this will pass notice by the rogue state, which may well conclude that it need not stare down the United States if it can scare the third party. Even if the third party would stand with the United States to present a united front against nuclear blackmail, it would have a harder time projecting the requisite confidence to make a stand, thereby emboldening the rogue state, which figures that its threats are more likely to crack the united front.

The United States may need options to keep the third party in line if it would leverage its cyberwar capabilities to bolster the allied position in a confrontation. The United States could, for instance, tell its ally to stand fast and play along with the U.S. assertion about the fragility of the rogue state's nuclear system, with an unspoken "or else" added for good measure. A softer option is to convince the third party that the *appearance* of steadfastness would dissuade the rogue nuclear state. Such argument makes equal sense, however, with or without a cyberwar counterthreat.

The third party's leadership may feel assured by U.S. confidence, but it may also consult its own military officers, who have enough daily contact with their U.S. counterparts to gauge exactly how much confidence U.S. commanders have in their ability to thwart the rogue's nuclear threat. Despite this, or perhaps because of this, the third party may want some indication that the United States could do what it suggests it can. What would constitute such evidence, and what would the United States feel comfortable showing the third party?

Any answer beyond "trust me on this" presumes a U.S. policy that is far more open about revealing offensive capabilities to other countries than it now seems to be.⁹ A crisis may

found and eradicated. This eliminates the possibility that the system was flawless. However, it says nothing about whether the system started out with one flaw or with two flaws—both were equally likely beforehand and equally likely afterward (if anything, finding one strengthens the case for there having been two flaws). So, once a flaw has been eradicated, there is *now* a 50-percent likelihood of there being zero remaining flaws and a 50-percent likelihood of there being one remaining flaw. The expected number of flaws after discovery and eradication is 0.5 (50 percent times 0 plus 50 percent times 1). Thus, in finding and fixing a flaw, while the actual number of flaws fell by one, the expected number of remaining flaws actually rose by 0.2.

⁹ An exception could be the "Five Eyes" consortium of the United States, United Kingdom, New Zealand, Australia, and Canada, with which we share detailed information on an ongoing basis.

create new risks associated with information-sharing. If standing fast requires pro-U.S. forces to project faith in the U.S. ability to nullify a nuclear threat, those nervous about taking such a huge risk, skeptics of cyberwar's power, or opponents of the United States within the allied government have every incentive to cast doubt on the proposition. They may be very tempted to leak selective information the United States has entrusted to them to argue that U.S. cyberwar capabilities are overstated. Penetration of third parties by agents of the nuclear rogue also cannot be dismissed.

Incidentally, a similar logic applies if the friendly third party is domestic (e.g., the U.S. Congress, opinion makers, the public). The more prominent a role cyberwar capabilities play—relative to retaliatory capabilities—in explaining why the United States is standing fast, the greater the demand to show why such confidence is warranted. It may be in the rogue state's interest, in fact, to argue that cyberwar capabilities are the *primary* basis for U.S. steadfastness, to pressure the United States to demonstrate what it can do.

So what *can* be shown? Representatives of the third-party government could be shown in real time how the rogue state's nuclear command system is operating. The hope is that such third parties accept that what they are shown is true (and that the target does not run honeynets to deceive cyber attackers). A better demonstration would tweak something in the rogue's nuclear command system that immediately shows up in something that the third party can independently observe and *then* hope they accept the fact that the exploit applies to all the rogue state's nuclear command systems.¹⁰ Thus, at a minimum, the United States would have to share very detailed knowledge about what it knows about the adversary's nuclear command system, from which, alas, sources and methods may be inferred.

The usual caveats apply. It may all be a bluff. The demonstration, although real, may fail. It may succeed and simply not be believed. The third party may feel that the rogue nuclear state has hidden workarounds. It helps if the United States understands the third party well enough to figure out what would surprise and delight it and then produce something that does exactly that.

And all this rests on the art of the possible in a world in which the rogue nuclear state has every incentive and every means to maintain command and control over its most important military asset.

Summation

The point of brandishing cyberattack capabilities in a nuclear crisis is to bolster U.S. confidence in not backing down before a grave threat. The more the United States demonstrates its belief that it could hack into the command system of a rouge nuclear state, the more likely it is that the United States can retain its freedom of action in a confrontation. The rogue state will understand that it must choose between restraint and annihilation, since it will see no possibility that the United States will yield.

The more the United States believes the rogue state will ignore even a credible cyberattack threat to its nuclear command-and-control system, the less value the United States would find

¹⁰ Familiar trade-offs come into play in weighing whether such observations should be hidden from the rogue state. If the rogue state finds out that something unexpected is happening in its nuclear command-and-control cycle, it may be motivated to attend to possible flaws but may also be warned against maneuvering itself into a position to be embarrassed. However, revealing the tweak that caused the effect is unhelpful because it would make fixing flaws easier without contributing much to the credibility of the cyberwar counterthreat.

in making such a cyberattack threat credible (since being explicit about how the cyberattack would work would make it easier to counter). The better the United States hides the details of its cyberwar capabilities, the less likely the rogue state is to reverse-engineer the hints and determine exactly what the exploit does. Thus, the better to frustrate the rogue nuclear state should it use nuclear weapons.

Conversely, the rogue nuclear state could read “hints” of U.S. capability as a scare tactic, something the United States would not employ if it actually had and thought it would use the capability. The nuclear rogue state might also conclude that the United States had so much (merited) confidence in its own capability that it believed “hints” would suffice to dissuade the rogue nuclear state.

A secondary benefit may come from injecting doubt into the rogue nuclear state *prior* to crisis. This may inhibit it from going so far down the crisis path that it cannot back down without a major loss of face and not cross the line at which the (imagined low) likelihood of avoiding retaliation is overwhelmed by the certainty of major embarrassment.

But what if success at a nuclear standoff requires the cooperation of a third-party state, which is typically at far greater risk from a rogue state’s nuclear weapons than the United States is? If promises and veiled threats cannot keep the third party in line, the latter may demand that the United States prove as much. Yet such proof may be unavailable short of a demonstration, which, if it fails, will discredit all other sources of confidence. If it succeeds, it may well tip off the nuclear rogue state, which then fixes the flaw that permitted the exploit. Furthermore, the third party may leak such information, particularly if there are people in or associated with the third-party state who want to see the U.S. argument discredited.

Nevertheless, there is something brittle about the United States relying too heavily on cyberattack capabilities to bolster its refusal to yield to a nuclear threat. An important, perhaps necessary, component of that tactic is the rogue state’s inability to discover *how* that would be so (since understanding its particulars can lead to nullifying such a capability). But the rogue state must still understand that it could be so if there is to be any possibility that brandishing cyberattack capabilities can help the United States manage the crisis. In short, the rogue nuclear state must credit what it cannot see, based on the reputation of the U.S. military.

Conclusions

Brandishing a cyber capability would do three things: declare a capability, suggest the possibility of its use in a particular circumstance, and indicate that such use would really hurt. In the era of the U.S.-Soviet nuclear standoff, the suggestion of use was the most relevant. Possession was obvious, and its consequences were well understood. The same does not hold true for cyberweapons. Possession is likely not obvious, and the ability to inflict serious harm is debatable. Even if demonstrated, what worked yesterday may not work today. But difficult does not mean impossible.

Advertising cyberwar capabilities may be helpful. It may back up a deterrence strategy. It might dissuade other states from conventional mischief or even from investing in mischief-making capabilities. It may reduce the other side's confidence in the reliability of its information, command-and-control, or weapon systems. In a nuclear confrontation, it may help build the edge that persuades other states that the brandisher will stay the course, thereby persuading the other states to yield.

Yet proving such capability is not easy, even if it exists. Cyber capabilities exist only in relationship to a specific target, which must be scoped to be understood. Cyber warriors can illustrate their ability to penetrate systems, but penetration is not the same as getting them to fail in useful ways. Since cyberattacks are essentially single-use weapons, they are diminished in the showing. It can be hard to persuade your friends that you have such capabilities when skepticism is in their interest.

Furthermore, brandishing may backfire. Touting an ability to strike back in cyberspace may communicate a tendency to shy from violence. Claiming the power to alter reality may convince others to blame the claimant when reality is disagreeable. Interfering with others' command and control may allow them to justify rules of engagement that abdicate their own responsibility over subordinates. And asserting an ability to nullify opposing nuclear systems may spur them to call what they perceive as a bluff.

Should the United States put the world on notice that it has cyber capabilities and knows how to use them? The wisdom of that course is not obvious. Evidence is scant that others act because they do not believe the United States has or can develop cyber capabilities. Conversely, the gains from brandishing such capabilities depend on the context and can be problematic even then.

There is both promise and risk in cyber brandishing, in both the conventional and nuclear cases. It would not hurt to give serious thought to ways in which the United States can enhance its ability to leverage what others believe are national capabilities. Stuxnet has certainly convinced many others that the United States can do many sophisticated things in cyberspace (regardless of what, if anything, the United States actually contributed to Stuxnet). This effort

will take considerable analysis and imagination, inasmuch as none of the various options presented here are obvious winners. That said, brandishing is an option that may also not work. It is no panacea, and it is unlikely to make a deterrence posture succeed if the other elements of deterrence (e.g., the will to wage war or, for red lines drawn in cyberspace, the ability to attribute) are weak.

References

- Alexander, Keith, "Advance Questions for Lieutenant General Keith Alexander, USA, Nominee for Commander, United States Cyber Command," statement to the U.S. Senate Committee on Armed Services, April 15, 2010, p. 21. As of June 29, 2011:
<http://www.armed-services.senate.gov/statemnt/2010/04%20April/Alexander%2004-15-10.pdf>
- Byman, Daniel L., Matthew C. Waxman, and Eric Larson, *Air Power as a Coercive Instrument*, Santa Monica, Calif.: RAND Corporation, MR-1061-AF, 1999. As of January 29, 2013:
http://www.rand.org/pubs/monograph_reports/MR1061.html
- Cliff, Roger, Mark Burles, Michael S. Chase, Derek Eaton, and Kevin L. Pollpeter, *Entering the Dragon's Lair: Chinese Antiaccess Strategies and Their Implications for the United States*, Santa Monica, Calif.: RAND Corporation, MG-524-AF, 2007. As of January 28, 2013:
<http://www.rand.org/pubs/monographs/MG524.html>
- "Cyber Strikes a 'Civilized' Option: Britain," *Technology Inquirer*, Agence France-Presse, June 3, 2012. As of June 3, 2012:
<http://technology.inquirer.net/11747/cyber-strikes-a-civilized-option-britain>
- Defense Advanced Research Projects Agency, "Cyber Experts Engage on DARPA's Plan X," press release, October 17, 2012. As of February 19, 2013:
<http://www.darpa.mil/NewsEvents/Releases/2012/10/17.aspx>
- Department of Defense, Office of General Counsel, "An Assessment of Legal Issues in Information Operations," May 1999.
- Fravel, M. Taylor, and Evan S. Medeiros, "China's Search for Assured Retaliation: The Evolution of Chinese Nuclear Strategy and Force Structure," *International Security*, Vol. 35, No. 2, Fall 2010, pp. 48–87.
- Goldman, Jeff, "Kaspersky Seeks Help Decrypting Gauss Malware Payload," *eSecurity Planet*, August 15, 2012. As of August 25, 2012:
<http://www.esecurityplanet.com/malware/kaspersky-seeks-help-decrypting-gauss-malware-payload.html>
- Harknett, Richard J., John P. Callaghan, and Rudi Kauffman, "Leaving Deterrence Behind: War-Fighting and National Cybersecurity," *Journal of Homeland Security and Emergency Management*, Vol. 7, No. 1, November 11, 2010.
- "Iran to Unveil National OS Soon," *PressTV*, January 4, 2011. As of June 3, 2012:
<http://www.presstv.ir/detail/158534.html>
- Johnson, David E., Karl P. Mueller, and William H. Taft, *Conventional Coercion Across the Spectrum of Operations: The Utility of U.S. Military Forces in the Emerging Security Environment*, Santa Monica, Calif.: RAND Corporation, MR-1494-A, 2003. As of January 28, 2013:
http://www.rand.org/pubs/monograph_reports/MR1494.html
- Libicki, Martin C., *Cyberdeterrence and Cyberwar*, Santa Monica, Calif.: RAND Corporation, MG-877-AF, 2009. As of January 28, 2013:
<http://www.rand.org/pubs/monographs/MG877.html>
- , "Wringing Deterrence from Cyberwar Capabilities," in Richmond M. Lloyd, ed., *Economics and Security: Resourcing National Priorities*, proceedings of a workshop sponsored by the William B. Ruger Chair of National Security Economics, Newport, R.I.: Naval War College, May 19–21, 2010, pp. 259–272.

Masters, Jonathan, "Confronting the Cyber Threat," New York: Council on Foreign Relations, May 23, 2011. As of February 4, 2013:

<http://www.cfr.org/technology-and-foreign-policy/confronting-cyber-threat/p15577>

Mearsheimer, John J., *Conventional Deterrence*, Ithaca, N.Y.: Cornell University Press, 1985.

Morgan, Forrest E., Karl P. Mueller, Evan S. Medeiros, Kevin L. Pollpeter, and Roger Cliff, *Dangerous Thresholds: Managing Escalation in the 21st Century*, Santa Monica, Calif.: RAND Corporation, MG-614-AF, 2008. As of January 28, 2013:

<http://www.rand.org/pubs/monographs/MG614.html>

Mueller, Karl P., Jasen J. Castillo, Forrest E. Morgan, Negeen Pegahi, and Brian Rosen, *Striking First: Preemptive and Preventive Attack in U.S. National Security Policy*, Santa Monica, Calif.: RAND Corporation, MG-403-AF, 2006. As of January 28, 2013:

<http://www.rand.org/pubs/monographs/MG403.html>

Nakashima, Ellen, "With Plan X, Pentagon Seeks to Spread U.S. Military Might to Cyberspace," *Washington Post*, May 30, 2012, p. A1.

Nowak, Martin. A., Karen M. Page, and Karl Sigmund, "Fairness Versus Reason in the Ultimatum Game," *Science*, Vol. 289, No. 5485, September 2000, pp. 1773–1775.

Pape, Robert A., *Bombing to Win: Air Power and Coercion in War*, Ithaca, N.Y.: Cornell University Press, 1996.

Quester, George H., *Deterrence Before Hiroshima*, Piscataway, N.J.: Transaction Publishers, 1986.

Rhoads, Christopher, and Farnaz Fassihi, "Iran Vows to Unplug Internet," *Wall Street Journal*, May 28, 2011. As of June 3, 2012:

<http://online.wsj.com/article/SB10001424052748704889404576277391449002016.html>

Sanger, David E., "Obama Order Sped Up Wave of Cyberattacks Against Iran," *New York Times*, June 1, 2012, p. 1.

Schelling, Thomas C., *Arms and Influence*, New Haven, Conn.: Yale University Press, 1966.

Schmidt, Marco F. H., and Jessica A. Sommerville, "Fairness Expectations and Altruistic Sharing in 15-Month-Old Human Infants," *PLOS ONE*, Vol. 6, No. 10, October 7, 2011.

Sterner, Eric, "Stuxnet and the Pentagon's Cyber Strategy," Arlington, Va.: George C. Marshall Institute, October 13, 2010. As of January 2013:

<http://www.marshall.org/article.php?id=918>

Zetter, Kim, "Senior Defense Official Caught Hedging on U.S. Involvement in Stuxnet," *WIRED*, May 26, 2011. As of May 27, 2012:

<http://www.wired.com/threatlevel/2011/05/defense-department-stuxnet/>