

Engineering the Cyber Defense Surge

OpenAI's Call for Collective Action Through the Lens of NIST SP 800-160

Dr. Ron Ross

CEO, RONROSSECURE, LLC.
ISTS Fellow, Dartmouth College

Abstract

OpenAI's open letter, "A Call for Collective Action on Cyber Defense," is timely and important. It warns that AI-enabled cyber-attacks are likely to become more widespread and sophisticated, argues that status quo security is inadequate, and calls for urgent action by organizations, technology providers, governments, and frontier AI companies. Its recommendations—least privilege, defense in depth, stronger access control, continuous testing, observability, monitoring, verified fixes, containment, recovery, and coordinated action—are directionally sound. This article makes a complementary argument: much of the engineering foundation needed to answer that call already exists. NIST Special Publication 800-160, Volume 1, Engineering Trustworthy Secure Systems, and Volume 2, Developing Cyber-Resilient Systems, established years ago a life-cycle, systems-engineering approach for building trustworthy secure and cyber-resilient systems. The central problem is therefore not the absence of a conceptual foundation. It is the persistent implementation gap between what the cybersecurity community has known how to engineer and what organizations have actually chosen, funded, acquired, built, operated, and sustained. AI changes the urgency, speed, scale, and economics of cyber conflict. It does not invalidate the foundational engineering principles. On the contrary, it makes them more consequential. The present moment should not be treated as a reason to discard prior systems security engineering work, but as a significant opportunity to finally operationalize it at scale and to extend it where AI creates genuinely new requirements.

1. A Timely Call and an Older Engineering Answer

OpenAI's call for collective cyber defense begins from a premise that should command broad agreement: the defensive status quo is not sufficient for the threat environment now emerging. The letter points to longstanding bugs, excessive permissions, misconfigurations, insecure and unpatched software, weak authentication, technical debt, and under-resourced defenders. It urges organizations to raise security standards with incident-level urgency and specifically calls for least privilege, strong access controls, defense in depth, continuous testing, verified fixes, observability, monitoring, containment, response, recovery, and collective action [1].

That call deserves support. The question is what *engineering model* should organize the surge. If the response becomes primarily a faster cycle of patching, tool deployment, AI-assisted detection, and incident response, defenders may improve operational tempo without correcting the structural conditions that make successful intrusion so consequential. AI can help defenders find and fix weaknesses faster, but it can also help adversaries find and exploit them faster. A durable defensive advantage therefore cannot depend only on winning a race to discover the next vulnerability.

NIST SP 800-160 offers a deeper answer. Volume 1 was first published in 2016 and substantially revised in 2022. It treats security as an *emergent property* of a system and applies systems security engineering across the complete system life cycle. Its stated purpose is to establish principles, concepts, activities, and tasks for engineering trustworthy secure systems of any purpose, type, scope, size, or complexity [2]. Volume 2, published in 2021, adds the complementary discipline of cyber resiliency engineering including architecting and sustaining systems capable of anticipating, withstanding, recovering from, and adapting to adverse conditions, attacks, and compromises [3].

Seen from that perspective, the OpenAI letter is not so much a repudiation of past cybersecurity thinking as a forceful contemporary restatement of needs that NIST SP 800-160 has been addressing for years. The important distinction is that OpenAI frames an urgent mobilization agenda, while SP 800-160 supplies much of the engineering discipline needed to make that mobilization durable.

The missing ingredient has not been another list of security controls. It has been the sustained engineering discipline to make security an emergent property of the system.

2. Where the OpenAI Recommendations and SP 800-160 Converge

OpenAI’s recommendations can be mapped directly to concepts in both volumes of SP 800-160. This is not evidence that the OpenAI letter is derivative; rather, it shows that two different communities, responding to the same worsening threat landscape and risk, are converging on many of the same necessities.

Table 1. OpenAI Recommendations and SP 800-160 Principles Converge

OpenAI Emphasis	SP 800-160 Foundation	Engineering Significance
Status quo security is insufficient	Adequate security; loss-driven engineering; commensurate protection and rigor	Security must be justified against consequences, not inherited from yesterday’s baseline.
Least privilege and strong access control	<i>Least Privilege; Mediated Access; Distributed Privilege; Minimal Trusted Elements</i>	Compromise should yield bounded authority , not system-wide power.
Defense in depth	<i>Defense in Depth; Domain Separation; Hierarchical Protection; Least Sharing</i>	Multiple independent barriers impede propagation after initial access.
Verify fixes and test continuously	<i>Substantiated Trustworthiness; Compositional Trustworthiness; Verification and Validation</i>	Security claims require evidence , not configuration assertions.
Contain attacks quickly	<i>Domain Separation; Protective Failure; Anomaly Detection; Segmentation</i>	Architecture should keep local compromise from becoming mission compromise .
Monitoring and observability	<i>Anomaly Detection; Analytic Monitoring; Contextual Awareness</i>	Detection must support effective, timely response and damage assessment .
Preparedness, response, and recovery	<i>Anticipate, Withstand, Recover, Adapt; Protective Recovery</i>	Resilience is engineered before the incident, not improvised during it.
Raise the bar for what is bought, built, and deployed	Life-cycle systems security engineering; supply-chain and enabling-system considerations	Acquisition and integration decisions create or remove trust dependencies .

The convergence is especially visible in the OpenAI letter’s recommendation to upgrade or replace systems to build in least privilege and defense in depth. SP 800-160, Volume 1, does not treat those ideas as isolated controls. They are part of a larger set of thirty principles for trustworthy secure design, including *Domain Separation, Least Sharing, Least Functionality, Mediated Access, Distributed Privilege, Protective Defaults, Protective Failure, Protective Recovery, Continuous Protection, Compositional Trustworthiness, Reduced Complexity, and Substantiated Trustworthiness* [2]. The principles are intended to be composed, tailored, revisited, and applied throughout engineering, not checked off after implementation.

That difference matters. Least privilege is more powerful when paired with domain separation, mediated access, minimal trusted elements, and distributed privilege. Defense in depth is more meaningful when the layers are sufficiently independent that one common-mode failure does not defeat all of them. Verification is stronger when it traces from stakeholder protection needs comprehensively through requirements, architecture, implementation, integration, operation, recovery, and disposal. The system, not the individual safeguard, becomes the unit of analysis.

The AI era has not made systems security engineering obsolete. It has made the cost of failing to practice it impossible to ignore.

3. The Critical Difference: Operational Cyber Defense Versus Systems Security Engineering

OpenAI's letter is intentionally action-oriented. It is aimed at a broad coalition and emphasizes immediate defensive priorities: fix high-risk weaknesses, deploy cyber-capable AI, test continuously, share threat intelligence, improve monitoring, support critical infrastructure, and coordinate response. Those are appropriate recommendations for a mobilization document [1].

SP 800-160 addresses a different layer of the problem. It asks how trustworthy secure behavior is engineered into the system in the first place. Volume 1 organizes the work around problem, solution, and trustworthiness contexts. The *problem* context establishes what must be protected and what losses must be prevented. The *solution* context develops the architecture and design that satisfy those needs. The *trustworthiness* context establishes justified confidence, through evidence, that the resulting system actually does what the security argument claims [2].

This distinction separates traditional cybersecurity activity from systems security engineering. Operational cybersecurity often begins with threats, vulnerabilities, alerts, and controls. Systems security engineering begins with the mission, stakeholder protection needs, unacceptable losses, system states, security constraints, architecture, and evidence. Both are necessary. But the second is what prevents the first from becoming an endless campaign of compensating for structural weaknesses that were designed into the system years earlier.

A five-alarm threat environment is the worst possible time to discover that the architecture was built on assumptions the adversary no longer has to respect.

4. SP 800-160 Was Designed for the World We Are Entering

One reason SP 800-160 is particularly relevant now is that it was deliberately *technology-agnostic*. The principles do not depend on a specific operating system, cloud architecture, malware family, or adversary toolchain. They reason about trust, authority, information flow, complexity, failure, recovery, composition, and evidence. Those concerns persist even as technologies and threats change.

The security design principles in Volume 1 explicitly seek to reduce uncertainty associated with adverse events by eliminating hazards, susceptibility, and vulnerability where practicable and controlling their effects where elimination is not possible [2]. That formulation is well suited to an AI-accelerated threat environment. AI may reduce the cost of reconnaissance, vulnerability discovery, exploit development,

social engineering, and attack-chain adaptation. Engineering cannot assume those capabilities will be eliminated. It can, however, reduce what an adversary gains from them.

A domain-separated architecture can require an attacker to repeatedly cross independently protected boundaries. Least privilege limits the authority inherited through a compromised identity or component. Least sharing reduces common services through which compromise propagates. Least functionality removes latent capabilities that become adversary tools after compromise. Mediated access constrains necessary interactions. Distributed privilege can make one compromise insufficient for consequential action. Protective failure can move a system toward a safer state when continued operation would increase loss. Protective recovery can require restoration from an authenticated trustworthy basis rather than merely restoring availability.

These are not merely defensive techniques. Together they alter the economics of attack. If AI makes initial compromise cheaper, architecture must make the conversion of that foothold into *mission consequence* more expensive. That is the durable contest.

AI can collapse the cost of finding a foothold. Architecture must raise the cost of turning that foothold into unacceptable loss.

5. Assume Compromise Before the Crisis Forces You To

The strongest connection between today's threat discussion and earlier NIST work may be SP 800-160, Volume 2. Its cyber-resiliency objective is not simply to prevent attacks. It is to engineer systems that can *anticipate, withstand, recover from, and adapt* to adverse conditions, stresses, attacks, and compromises [3]. The framework directly addresses the limitation of *prevention-centered security*. A system may be penetrated, a component may fail, a supplier may be compromised, an administrator credential may be stolen, or an AI-enabled adversary may discover a previously unknown exploit. The engineering question then becomes: **what does the system permit that condition to become?**

SP 800-160, Volume 2's strategic design principle *Assume Compromised Resources* is particularly important. It directs engineers to recognize that system resources can be compromised for extended periods without detection and to design accordingly. The principle of *Expect Adversaries to Evolve* makes a related point: **systems should not be engineered only against adversary behavior already observed**. The combination anticipates precisely the environment now motivating calls for urgent collective defense.

Volume 2 also provides fourteen cyber-resiliency *techniques*, including segmentation, non-persistence, substantiated integrity, adaptive response, analytic monitoring, deception, diversity, dynamic positioning, redundancy, privilege restriction, coordinated protection, contextual awareness, and others [3]. These constructs go beyond the idea of adding more security products. They shape how the system behaves when adversity succeeds.

That combination of principles gives SP 800-160, Volume 2 direct application to a threat environment defined by faster reconnaissance and faster adaptation, not because the framework was written with AI in mind, but because it was written for adversaries who evolve, which AI now allows them to do at a different tempo.

6. From Initial Compromise to Mission Loss: The Gap Architecture Must Control

This distinction is developed further in the companion analysis *Engineering for Compromise* [4], which separates component compromise, system compromise, hazardous system state, and unacceptable mission loss as four distinct events. A successful penetration creates an adversary condition. Whether that condition becomes catastrophe depends on the architecture.

The practical question is severe but useful: if this account, component, service, supplier, maintenance tool, AI agent, or administrative domain were completely controlled by an adversary, what could it cause? What information could it reach? What authority could it exercise? What interfaces could it invoke? What states could it change? What shared dependencies could it exploit? What other domains could it reach? The answer exposes the true consequence of lost trust.

This is where the design principles in SP 800-160 can deepen the OpenAI call. Continuous testing, AI-assisted vulnerability discovery, rapid patching, and stronger monitoring can reduce the probability and duration of compromise. Architecture determines the blast radius. The objective should be to break the *causal chain* before a local failure becomes a system-wide or mission-level loss.

That framework uses compromise containment boundaries, propagation analysis, and evidence to ask whether the system can preserve critical properties even when selected elements become untrustworthy. This is not a replacement for prevention. It is what comes next when prevention is treated as necessary but fallible [4].

The security objective should not be ‘nothing is ever compromised.’ It should be ‘compromise does not automatically become catastrophe.’

7. Verification: ‘Fix It’ Is Not Complete Until the Claim Is Substantiated

OpenAI repeatedly emphasizes verified fixes, continuous testing, observability, and measuring whether defenses actually work [1]. That is one of the most important points in the letter, and it aligns strongly with the *trustworthiness* context in SP 800-160.

Security programs too often equate implementation with effectiveness. A security control is deployed, a configuration is changed, a vulnerability is patched, or a policy is issued, and the activity is treated as evidence of security. SP 800-160 demands a stronger argument. Trustworthiness is a judgment supported by *evidence*. Demonstration, inspection, analysis, testing, penetration testing, adversarial evaluation, fault injection, recovery exercises, and other methods can contribute evidence, but no single method establishes the trustworthiness of the whole system.

Compositional Trustworthiness is especially relevant. Individually trustworthy components do not automatically compose into a trustworthy system. Integration can create new shared dependencies, privilege paths, bypasses, timing relationships, recovery assumptions, and common-mode failures. Likewise, a verified patch can close one vulnerability without demonstrating that the mission-level security constraint is preserved.

The engineering goal should therefore be an evidence chain: **Mission → Unacceptable Loss → Hazardous System Condition → Security Constraint → Architecture → Mechanism → Evidence** [4]. This creates traceability from what matters to why a particular safeguard exists and how confidence in its effectiveness is substantiated.

In an AI era, that evidence must also have an operational lifetime. A model update, new agentic capability, changed dependency, revised threat assumption, cloud migration, emergency maintenance path, or supply-chain substitution can invalidate earlier assurance. Continuous monitoring is therefore not merely alert collection; it is part of maintaining the validity of the *security argument*.

8. The Implementation Gap: Why Sound Engineering Did Not Become the Default

It is tempting to describe the history simply as complacency. There is truth in that diagnosis, but it is incomplete. The more useful conclusion is that the cybersecurity ecosystem developed powerful incentives around operational controls, compliance, vulnerability management, product acquisition, and incident response, while the harder work of engineering security into system architecture often remained secondary.

Several forces contributed. Legacy systems constrained architectural change. Security organizations were frequently separated from systems engineering and acquisition. Compliance regimes made controls visible and auditable, while architectural qualities were harder to measure. Procurement cycles rewarded features and delivery schedules. Organizations accepted technical debt because the near-term business case for redesign was difficult to quantify. Security teams were under-resourced, as OpenAI itself observes. And when major attacks remained episodic, organizations could rationally, if shortsightedly, defer expensive structural changes. The result is an implementation gap. The community accumulated extensive guidance on least privilege, separation, minimality, resilience, recovery, assurance, and life-cycle engineering, yet many deployed systems remained highly connected, privilege-rich, dependency-heavy, difficult to analyze, and operationally brittle.

A fair skeptic will raise the obvious objection here: if this engineering foundation has been publicly available for the better part of a decade and adoption still fell short, why should this moment produce a different outcome? It is a reasonable question, and the honest answer is not that organizations have suddenly become more virtuous. Nothing about SP 800-160 itself has changed. What has changed is the cost of deferral. AI-enabled offense compresses time, reduces attacker cost, increases scale, and may expose accumulated technical debt faster than institutions can compensate operationally. That shift does not guarantee organizations will act on guidance they previously deprioritized. It does mean that continuing not to act now carries a materially different order of risk than it did five years ago, and that the business case for deferral, which was always debatable, is becoming harder to make.

The appropriate lesson, then, is not to assign blame to organizations that failed to perfectly implement a demanding engineering discipline. It is to recognize that the *economics* have changed enough that continued deferral is no longer a plausible strategy for organizations operating mission critical systems.

AI-enabled offense compresses time, reduces attacker cost, increases scale, and may expose accumulated technical debt faster than institutions can compensate operationally.

9. Comparing OpenAI Recommendations to SP 800-160 Engineering Guidance

A fair comparison must also identify *differences*. OpenAI's letter adds contemporary priorities that SP 800-160 was not designed to prescribe in detail. It calls for broad access to cyber-capable AI, lower-cost defensive models, frontier capabilities for hard problems, responsible model access, agentic identity traceability, AI-powered security tools, threat-intelligence sharing, funding for under-resourced critical infrastructure, and a global coalition of technology providers, governments, and defenders [1]. Those are

ecosystem, policy, and technology-deployment recommendations rather than systems-engineering constructs.

Related work extending this framework to frontier-AI-specific threats reaches a compatible conclusion [5]: SP 800-160 remains the architectural backbone, while AI-specific threat modeling, data and model integrity, machine-speed defensive adaptation, and governance considerations must be layered on top of it. AI-generated code should be treated as untrusted input until validated; model and data provenance matter; agentic authority must be bounded; and defensive tempo must increasingly operate at machine speed.

Conversely, SP 800-160 goes substantially deeper than the OpenAI letter in areas that a short mobilization document cannot. It provides a system life cycle engineering framework; a concept of protection needs, asset loss, and adequate security; thirty design principles for trustworthy secure systems; explicit treatment of trustworthiness and assurance; problem, solution, and trustworthiness contexts; engineering processes from mission analysis through disposal; and, in Volume 2, a structured cyber-resiliency framework of goals, objectives, techniques, approaches, and strategic design principles [2][3].

The two bodies of work should therefore be viewed as complementary. OpenAI supplies urgency, coalition, AI-era resources, and a call to mobilize. SP 800-160 supplies a mature engineering foundation for deciding what to build, how to structure it, how it should behave under adversity, and what evidence is needed to justify confidence.

10. A Practical Synthesis for the Collective Cyber Defense Surge

The strongest response would combine the urgency of the OpenAI letter with the *engineering discipline* of SP 800-160. That synthesis can be expressed as a sequence of actions.

- 1. Start with mission and unacceptable loss, not the tool catalog.** Identify the services, capabilities, data, control functions, and societal outcomes that cannot be allowed to fail. Translate those consequences into explicit protection needs and security constraints.
- 2. Assume selected resources will be compromised.** For critical systems, analyze the consequences of hostile control of identities, endpoints, cloud services, suppliers, management planes, AI agents, recovery systems, and maintenance paths. Design containment before the incident.
- 3. Re-architect privilege and connectivity.** Use least privilege, domain separation, least sharing, least functionality, mediated access, distributed privilege, and minimal trusted elements to reduce the authority and reachability an adversary acquires from any one foothold.
- 4. Engineer resilience, not merely resistance.** Use the anticipate-withstand-recover-adapt model and cyber-resiliency techniques to preserve essential mission functions under adversity, including when prevention and detection are imperfect.
- 5. Use AI to accelerate the defensive work.** Apply AI to vulnerability discovery, code review, configuration analysis, monitoring, triage, testing, and response—but place AI components inside architectures that independently constrain their authority, access, autonomy, and potential impact.
- 6. Substantiate security claims with evidence.** Verify that fixes work, test boundary behavior, exercise recovery, challenge assumptions, red-team the integrated system, and maintain assurance evidence as the system and threat environment change.
- 7. Carry security intent through the entire life cycle.** Mission changes, acquisitions, integrations, updates, maintenance, emergency procedures, and disposal can all invalidate earlier assumptions. Security must remain an engineering thread, not an operational phase.

8. Measure mission consequence, not just cyber activity. Track whether attacks are contained, whether critical functions survive, whether recovery restores trustworthy state, and whether compromise margins remain sufficient—not only vulnerabilities closed or alerts processed.

This synthesis also changes how organizations should interpret the current “defenders’ window.” The window should not be used only to patch the backlog before AI makes exploitation faster. It should be used to remove the architectural conditions that allow one successful exploit to produce disproportionate consequences. Otherwise, the community may emerge from the surge with better tools but essentially the same fragile systems.

Use the defenders’ window not only to fix vulnerabilities, but to fix the architecture that makes vulnerabilities catastrophic.

11. National Leadership as a Forcing Function

The preceding sections describe an *engineering* problem. This section addresses a narrower, adjacent question: how adoption might actually accelerate, given that the implementation gap described above has persisted for years despite guidance that was already available.

OpenAI frames cyber defense as a collective-action problem, and that framing carries a policy implication worth stating directly. Collective awareness may not be sufficient when the costs of architectural modernization are borne by individual organizations while the consequences of systemic weakness are national.

For the nation's most consequential systems, national leadership is one plausible *forcing function*, through federal legislation, executive orders, OMB policy, directives, acquisition requirements, or sector-specific regulatory action, where appropriate. The objective would not be to prescribe a single architecture from the federal government. It would be to establish a clear national expectation that systems supporting critical missions and essential services be engineered to remain trustworthy secure and resilient against the adversaries they are expected to face [6].

This is not an argument for indiscriminate cybersecurity regulation. It is, however, an observation that federal leadership can supply the priorities, accountability, and acquisition signals needed to turn SP 800-160 from respected guidance into sustained practice, something market incentives and voluntary adoption have not yet accomplished at sufficient scale in the decade since SP 800-160 was first published.

Collective cyber defense may ultimately require more than collective awareness. It may require national leadership capable of turning sound engineering principles into sustained national priorities.

12. From a Five-Alarm Threat to a Long-Overdue Engineering Reset

OpenAI and its fellow signatories are right to call for urgency. AI-enabled cyber capabilities are changing the economics and tempo of offense, and critical infrastructure cannot depend on incremental improvement alone. The response should be broad, well-funded, internationally coordinated, and equipped with the best defensive AI capabilities available.

But urgency should not be confused with *novelty*. Many of the foundational ideas now being rediscovered under pressure—least privilege, defense in depth, separation, bounded trust, continuous protection, verified effectiveness, resilience, recovery, monitoring, and adaptation—have long been part of the systems security engineering discipline. NIST SP 800-160 assembled those ideas into a coherent engineering framework precisely because trustworthy secure systems cannot be achieved by security controls alone.

The most respectful way to read the OpenAI letter is therefore as an opportunity for convergence. The AI community is bringing extraordinary technical capability, investment, visibility, and urgency to cyber defense. The systems security engineering community brings decades of accumulated knowledge and expertise about how to build trustworthy secure systems and reason about their behavior under adversity. Neither is sufficient alone.

The challenge now is implementation. Organizations should not merely cite SP 800-160, map it to a framework, or treat its principles as aspirational language. They should use it to change architectures, acquisition decisions, trust relationships, privilege structures, failure behavior, recovery mechanisms, assurance evidence, and life-cycle processes. At the same time, SP 800-160 should be extended and operationalized for AI-specific threats where model behavior, data/model integrity, machine-speed attack and defense, and agentic autonomy create genuinely new engineering demands.

NIST SP 800-160 was ahead of its time in one important sense: it treated security as an *emergent system property* to be deliberately engineered across the system life cycle before the current threat environment made that position unavoidable. The five-alarm condition we face today does not diminish that work. It validates its central premise.

We do not need to invent the foundations of trustworthy cyber defense from scratch. We need to finally engineer them at the scale and urgency the threat now demands.

13. Conclusion

The OpenAI call for collective cyber defense should be welcomed. Its warning is credible, its emphasis on critical infrastructure is appropriate, and its recommendations for urgent remediation, AI-enabled defense, verification, monitoring, containment, recovery, and collective action are important [1].

The next step is to connect that mobilization to a mature engineering foundation. NIST SP 800-160, Volume 1, provides the discipline for engineering trustworthy secure systems from mission needs through architecture, implementation, verification, operation, maintenance, and disposal. Volume 2 provides the complementary discipline for engineering systems to anticipate, withstand, recover from, and adapt to adversity and compromise [2][3].

Together, they support a proposition that is more demanding than “improve cybersecurity”: engineer systems whose security properties are deliberate, compositional, evidence-based, resilient, and sustained throughout the life cycle. In an era of AI-accelerated attack, that is no longer an advanced aspiration reserved for a small class of high-assurance systems. For critical missions and essential services, it is becoming the minimum responsible engineering posture.

The cyber community has a rare opportunity. The urgency is new. Many of the engineering principles are not. The task is to close the gap between what we have known and what we have built.

Acknowledgments

The author gratefully acknowledges the support provided by Dartmouth College's Institute for Security, Technology, and Society (ISTS). The resources provided by the ISTS contributed to the research and development of this paper. The views and conclusions expressed herein are those of the author and do not necessarily represent the views of Dartmouth College.

References

- [1] OpenAI and supporting organizations, "A Call for Collective Action on Cyber Defense," OpenAI, 2026. <https://openai.com/collective-cyberdefense>
- [2] R. Ross, M. Winstead, and M. McEvilly, *Engineering Trustworthy Secure Systems*, NIST Special Publication 800-160, Volume 1, Revision 1, National Institute of Standards and Technology, November 2022. <https://doi.org/10.6028/NIST.SP.800-160v1r1>
- [3] R. Ross, V. Pillitteri, R. Graubart, D. Bodeau, and R. McQuaid, *Developing Cyber-Resilient Systems: A Systems Security Engineering Approach*, NIST Special Publication 800-160, Volume 2, Revision 1, National Institute of Standards and Technology, December 2021. <https://doi.org/10.6028/NIST.SP.800-160v2r1>
- [4] R. Ross, "Engineering for Compromise: Designing Systems That Remain Trustworthy Secure Under Adversity," RONROSSECURE, LLC / LinkedIn, August 2026.
- [5] R. Ross, "A Multidimensional Protection Strategy for Engineering Mission-Resilient Systems in the Age of Frontier AI," RONROSSECURE, LLC / LinkedIn, August 2026.
- [6] R. Ross, "An Executive Order for Trustworthy Systems?," RONROSSECURE, LLC / LinkedIn, June 2026.