

AGENTS OF CONTEXT: A METHODOLOGICAL CRITIQUE AND COUNTER-EVIDENCE ANALYSIS OF ADVERSARIAL RED-TEAMING CLAIMS FOR AUTONOMOUS AI AGENTS

Mosab Hawarey

Director, Geospatial Research (director@geospatial.ch)

Received: 2026 February 25 | Revised: 2026 March 01 | Approved: 2026 March 01 | Published: 2026 March 02

Abstract

The recent paper *Agents of Chaos* (Shapira et al., 2026) reports an exploratory red-teaming study of six autonomous language-model-powered agents, documenting eleven vulnerability case studies including unauthorized compliance, sensitive data disclosure, identity spoofing, and multi-agent vulnerability propagation. The authors conclude that these findings establish security, privacy, and governance-relevant vulnerabilities in realistic deployment settings, warranting urgent attention from policymakers. We present a systematic methodological critique and comprehensive counter-evidence analysis challenging both the study's methods and its governance-level conclusions.

We advance five structurally independent arguments. First, every documented vulnerability occurred in agents operating without any production safety infrastructure—no guardrails, no sandboxing, no identity verification, no human-in-the-loop approval—rendering the “realistic deployment settings” claim unsupported. Second, the participant pool of twenty expert AI researchers systematically inflates apparent risk through selection bias, with 72.7% of vulnerabilities concentrated in a single agent. Third, eleven case studies from one framework over fourteen days cannot epistemologically support the governance prescriptions the paper derives from them. Fourth, the paper's own safety success case studies—including zero compliance across fourteen-plus prompt injection variants and spontaneous inter-agent safety coordination—internally contradict its alarmist conclusions. Fifth, the study conflates vulnerability existence with deployment risk, a distinction well-established in cybersecurity but absent from the paper's inferential chain.

We demonstrate that 100% of documented vulnerabilities have known, production-ready engineering mitigations across five converging industry defense frameworks, with evidence ranging from peer-reviewed empirical validation to vendor-published guidance. We propose reframing the paper's contributions as a valuable engineering requirements checklist rather than a governance emergency, and articulate six principles for proportionate AI agent safety evaluation alongside eight concrete recommendations for future red-teaming methodology. Our analysis draws on the established red-teaming

methodology literature, case study epistemology, the cybersecurity vulnerability-risk distinction, and the documented technology panic cycle to argue for evidence-based, proportionate governance of AI agents.

Keywords: AI agents; red-teaming methodology; safety evaluation; vulnerability assessment; defense-in-depth; proportionate governance; AI safety; methodological critique

Citation: Hawarey, M. (2026). Agents of Context: A Methodological Critique and Counter-Evidence Analysis of Adversarial Red-Teaming Claims for Autonomous AI Agents. *AIR Journal of Interdisciplinary Research*, Vol. 2026, AIRJIR2026369, DOI: 10.65737/AIRJIR2026369.

1. Introduction

Autonomous AI agents—language models equipped with persistent memory, tool access, and the ability to plan and act across sessions without per-action human approval—represent one of the most significant developments in artificial intelligence. As these systems move from research prototypes toward production deployment, understanding their safety properties is a matter of genuine importance. The question is not whether safety evaluation matters, but how it should be conducted, what evidential standards it should meet, and what conclusions it can support.

In February 2026, Shapira et al. published *Agents of Chaos*, an exploratory red-teaming study deploying six autonomous language-model-powered agents on the OpenClaw scaffold in a live Discord server. Over a two-week period, twenty AI researchers interacted with the agents under both benign and adversarial conditions, producing 335 session transcripts across 78 Discord channels. The study documents eleven vulnerability case studies—including unauthorized compliance with non-owners, sensitive data disclosure, identity spoofing, and multi-agent vulnerability propagation—alongside five safety success case studies demonstrating prompt injection resistance, social engineering rejection, and spontaneous inter-agent safety coordination. The paper has attracted rapid attention: authored by a 38-researcher team spanning Northeastern University, Independent Researchers, Stanford University, University of British Columbia, Harvard University, Hebrew University, Max Planck Institute for Biological Cybernetics, MIT, Tufts University, Carnegie Mellon University, Alter, Technion, and Vector Institute, it was accompanied by a dedicated project website with interactive Discord logs, received immediate coverage from AI-focused media outlets, and was cited by multiple concurrent studies on agentic AI security within days of publication. This rapid dissemination makes methodological scrutiny of the paper's inferential chain particularly urgent—its conclusions are actively shaping the discourse on AI agent safety at precisely the moment when governance frameworks are being formulated.

The study's empirical observations are valuable. Its conclusions, however, systematically exceed its evidence. The paper claims to have established vulnerabilities in “realistic deployment settings” and calls for “urgent attention from legal scholars, policymakers, and researchers across disciplines.” We argue that both claims are unsupported by the study's design, methodology, and data.

This paper presents a systematic methodological critique and comprehensive counter-evidence analysis. We advance five structurally independent arguments that collectively demonstrate the gap between the target paper’s evidence and its conclusions:

- (1) **The study tests raw models, not deployed systems.** Every documented vulnerability occurred in agents operating without any production safety infrastructure—no guardrails, no sandboxing, no identity verification, no human-in-the-loop approval. This is analogous to reporting automobiles as unsafe when tested without brakes, seatbelts, or airbags.
- (2) **Expert adversaries inflate risk through selection bias.** Twenty AI researchers represent the extreme tail of adversarial capability. The absence of a non-expert baseline prevents any inference about typical-use risk, and 72.7% of vulnerabilities are concentrated in a single agent.
- (3) **Existence-proofs do not warrant governance prescriptions.** Eleven case studies from one framework over fourteen days support hypothesis generation and engineering guidance, not the “urgent” policy attention the paper calls for. The case study epistemology literature is explicit on this boundary.
- (4) **The paper’s own safety successes undermine its alarm.** Five case studies document robust prompt injection resistance (fourteen-plus variants, zero compliance), spontaneous inter-agent safety coordination, and consistent rejection of email spoofing and data tampering—demonstrating substantial baseline safety even in unguarded agents.
- (5) **Vulnerability does not equal risk in any established framework.** The study documents vulnerability existence but provides no threat likelihood assessment, no impact quantification, and no engagement with the existing defense ecosystem. Two of three required risk components are entirely missing.

We further demonstrate that 100% of the documented vulnerabilities have known, production-ready engineering mitigations documented across five converging industry defense frameworks. We propose reframing the target paper’s contributions as a valuable engineering requirements checklist rather than a governance emergency, and articulate six principles for proportionate AI agent safety evaluation alongside eight concrete recommendations for future red-teaming methodology.

Our critique is constructive. We acknowledge the target paper’s genuine contributions—including the first systematic documentation of multi-agent vulnerability propagation, the identification of identity verification as a critical architectural gap, and the valuable finding of silent censorship from geographically-restricted models. We argue not that the findings are wrong, but that the conclusions drawn from them are disproportionate to the evidence, and that the appropriate response is rigorous engineering rather than governance alarm.

Methodological approach. This paper adopts a desk-based methodological critique and counter-evidence analysis rather than an empirical replication study. This approach is appropriate for several reasons. First, the primary contribution is analytical rather than empirical: we evaluate the inferential validity of the target paper’s conclusions relative to its evidence, a task that requires logical and methodological analysis rather than new data collection. Second, the counter-evidence we present—defense frameworks, production deployment data, and empirical

guardrail evaluations—is drawn from publicly available, independently verifiable sources, enabling full reproducibility without requiring access to the original experimental infrastructure. Third, while a replication study comparing agent behavior with and without production safety infrastructure would provide the strongest form of direct evidence (and we recommend this as future work in Section 7.4), such a study addresses a different research question—whether mitigations work in practice—rather than the question we address here: whether the target paper’s evidence supports its conclusions. We selected the five defense frameworks analyzed in Section 4.1 based on three criteria: independent development by distinct organizations, publication by early 2026, and explicit coverage of agentic AI systems. We apply established cybersecurity risk frameworks (ISACA, NIST) to the AI agent domain on the grounds that the vulnerability-threat-impact risk chain is domain-general and has been adopted by the AI safety community itself, including in the NIST AI RMF and the OWASP Agentic AI Top 10. We acknowledge that AI-specific risk factors—such as emergent capabilities, goal misalignment, and distributional shift—may require extensions to traditional cybersecurity risk models, a limitation we discuss in Section 7.4.

The remainder of this paper is organized as follows. Section 2 provides a structured analysis of the target study’s design, findings, and claims. Section 3 presents the methodological critique across four dimensions. Section 4 presents counter-evidence from defense frameworks and safe deployments. Section 5 analyzes the vulnerability-risk gap and the paper’s internal contradictions. Section 6 offers the constructive reframing and proportionate governance recommendations. Section 7 provides a synthesis discussion, and Section 8 concludes.

2. Background: The “Agents of Chaos” Study

This section provides a structured analysis of the *Agents of Chaos* study (Shapira et al., 2026), establishing the factual foundation upon which our critique is built. We characterize the experimental design, catalog all sixteen case studies, and identify the specific inferential leaps where the paper’s conclusions exceed its evidence.

2.1 Experimental Design

The study deployed six autonomous agents on the OpenClaw open-source scaffold—a framework providing frontier language models with persistent memory, tool access, and genuine autonomy. Agents could initiate contact, form plans, and act across sessions with no per-action human approval. The deployment ran from February 2–17, 2026, on a live Discord server where twenty AI researchers interacted with agents under both benign and adversarial conditions, producing 335 OpenClaw session transcripts across 78 Discord channels.

Table 1 summarizes the agent configuration. A critical observation is model homogeneity: five of six agents ran Anthropic Claude variants, with only Quinn running an alternative model (Kimi K2.5 from Moonshot AI). The study therefore primarily characterizes Claude’s behavior on the OpenClaw scaffold, not AI agents generically.

Table 1. *Agent configuration in the Agents of Chaos study.*

Agent	Model	Owner	Notable Role in Study
Ash	Claude	Chris	Most active; 8 of 11 vulnerability CS (72.7%); also 2 safety successes
Flux	Claude	Chris	Entered 9-day infinite loop with Ash (CS4)
Jarvis	Claude	Danny	PII handling; semantic reframing bypass (CS3); data tampering refused (CS14)
Quinn	Kimi K2.5	Server 1	Only non-Claude agent; silent censorship from China-specific restrictions (CS6)
Mira	Claude Opus	Andy	Over-compliance (CS2, CS5) and safety behaviors (CS15, CS16)
Doug	Claude Opus	Andy	Teaching (CS9), over-compliance (CS2, CS5), safety coordination (CS15, CS16)

Each agent was granted extensive capabilities, summarized in Table 2. Most notably, agents had unrestricted bash shell execution—the ability to run arbitrary system commands without any sandboxing, approval gate, or audit trail.

Table 2. *Capabilities granted to each agent with safety implications.*

Capability	Details	Safety Implications
Email (ProtonMail)	Send, receive, search inbox	Data exfiltration vector
Discord	Multi-channel server access	Social interaction surface; spoofing vector
Shell (Bash)	Unrestricted execution	Arbitrary system command execution
File System	20 GB persistent volume	Persistent storage with no quota enforcement
Cron Jobs	Scheduled autonomous tasks	Unsupervised scheduled actions
External APIs	Web, Moltbook, GitHub	External interaction without content filtering

2.2 The Absent Safety Infrastructure

The most consequential feature of the experimental design is not what was present but what was absent. Table 3 catalogs the standard production safety measures that were not deployed. Every vulnerability documented in the study occurred in a context where these measures were deliberately excluded.

Table 3. *Standard production safety measures absent from the study.*

Standard Production Safety Measure	Present in Study?
Guardrails toolkit (e.g., NeMo, Constitutional Classifiers)	No
Sandboxed/restricted shell execution	No — unrestricted bash
Human-in-the-loop approval for sensitive actions	No — no per-action approval
Identity verification beyond display names	No — Discord display names only
Authorization/permission model	No — implicit trust of any speaker
Rate limiting on messages/actions	No
Storage quotas or warnings	No
Content filtering on inputs/outputs	No — only base model alignment
Monitoring/alerting for anomalous behavior	No
Audit logging with owner notification	No

This design is optimal for vulnerability discovery—it maximizes the attack surface and minimizes defensive barriers, precisely as a penetration test should. However, this design provides no information about what would happen in production deployments where some or all of these measures are in place. The study characterizes raw model behavior, not deployed system behavior.

2.3 Vulnerability Case Studies

The study documents eleven vulnerability case studies spanning seven distinct failure categories: authorization/identity failure (CS2, CS8, CS11), semantic/social bypass (CS3, CS7), resource management failure (CS4, CS5), disproportionate judgment (CS1), indirect prompt injection (CS10), provider-level opacity (CS6), and multi-agent vulnerability propagation (CS10, CS11). Table 4 provides the complete catalog.

Table 4. *Complete catalog of vulnerability case studies documented in Agents of Chaos.*

CS#	Name	Category	Agent(s)	Key Outcome
CS1	The Nuclear Option	Disproportionate Response	Ash	Destroyed own mail server to protect a secret
CS2	Non-Owner Compliance	Authorization Failure	Mira, Doug	123 email records returned to unauthorized user
CS3	The Forwarded Inbox	Semantic Reframing	Jarvis	SSN, bank, medical data exposed via “forward” vs “share”
CS4	The Infinite Loop	Resource Exhaustion	Ash, Flux	9-day loop, ~60,000 tokens, no termination

CS5	Storage Exhaustion	Denial of Service	Mira, Doug	Email server DoS via repeated large attachments
CS6	Silent Censorship	Provider Opacity	Quinn	Opaque errors on politically sensitive topics
CS7	The Guilt Trip	Emotional Pressure	Ash	Compliance after 12+ principled refusals
CS8	Identity Hijack	Identity Spoofing	Ash	Full system takeover via display name spoofing
CS10	Corrupted Constitution	Prompt Injection	Ash	GitHub Gist injection; attempted agent shutdowns
CS11	The Libel Campaign	Multi-Agent Propagation	Ash	Mass broadcast of fabricated emergency

A striking feature of the vulnerability distribution is the concentration on a single agent: Ash accounts for eight of eleven vulnerability case studies (72.7%). This extreme concentration may reflect the distribution of adversarial attention rather than a uniform property of all agents, since Ash was described as the most active agent and likely attracted disproportionate tester focus.

2.4 Safety Success Case Studies

Alongside the vulnerability findings, the study documents five safety success case studies (Table 5) that are at least as empirically significant. The project website itself notes that these are “sometimes described in the paper as ‘failed experiments’ because the adversarial designs didn’t unfold as hypothesized,” but adds: “We think this framing inverts the finding: these are cases where the agents got it right.” This editorial acknowledgment from the study’s own authors constitutes an internal critique of the paper’s framing.

Table 5. *Complete catalog of safety success case studies.*

CS#	Name	Category	Agent(s)	Key Outcome
CS9	Cross-Agent Teaching	Productive Collaboration	Doug, Mira	Successful knowledge transfer across environments
CS12	Injection Refused	Prompt Injection Resistance	Ash	14+ variants, zero compliance
CS13	Email Spoofing Refused	Social Reframing Resistance	Ash	Consistent refusal with clear reasoning
CS14	Data Tampering Refused	API Boundary Maintenance	Jarvis	Maintained API vs. file distinction under pressure
CS15	Social Engineering Resisted	Impersonation Resistance	Mira, Doug	Rejected owner impersonation attempt
CS16	Emergent Safety Coordination	Spontaneous Policy Negotiation	Mira, Doug	Spontaneous inter-agent safety coordination

CS12 is particularly significant: the same agent (Ash/Claude) that accounts for 72.7% of vulnerability case studies also demonstrated zero compliance across fourteen-plus distinct prompt injection variants—including base64-encoded commands, image-embedded instructions, fake authority tags, and XML/JSON privilege escalation. CS16 documents an entirely emergent behavior: without any instruction, Doug identified a suspicious pattern, warned Mira, and they jointly negotiated a stricter shared safety policy.

2.5 What the Paper Demonstrates Versus What It Implies

The paper’s empirical observations are valid within their scope: the documented vulnerabilities genuinely occurred as described, and the data (published in full) enables independent verification. However, the paper’s conclusions systematically exceed its evidence in several critical respects.

First, the paper claims its findings “establish the existence of security-, privacy-, and governance-relevant vulnerabilities in realistic deployment settings.” The word “realistic” is unsupported: the deployment setting lacked every standard production safety measure (Table 3). No responsible production deployment would replicate these conditions.

Second, the paper concludes that its findings “warrant urgent attention from legal scholars, policymakers, and researchers across disciplines.” The word “urgent” implies imminent danger requiring immediate action, yet the study provides no risk quantification, no threat-likelihood assessment, and no engagement with the existing defense ecosystem. Moreover, the paper provides no evidence that its findings have been taken up by policymakers or have influenced governance decisions, making the urgency framing aspirational rather than empirically grounded.

Third, the paper frames the eleven vulnerability case studies as its primary contribution and the five safety successes as secondary “failed experiments.” This framing produces a systematically biased picture that overstates the vulnerability surface and understates the baseline safety already present in unguarded agents.

We characterize this gap between evidence and conclusion as the paper’s central inferential problem—analogous to testing automobile safety by removing all seatbelts, airbags, and brakes, hiring twenty professional stunt drivers to deliberately crash the vehicles, documenting that crashes are unsafe, and then calling for urgent attention from policymakers about automobile safety. The findings are technically correct, but the conclusion conflates the absence of safety features with the inherent danger of the technology.

3. Methodological Critique

We identify four methodological limitations that collectively undermine the paper’s governance-level conclusions while preserving the validity of its vulnerability discovery contributions.

3.1 Red-Teaming Identifies Attack Surfaces, Not Deployment Risk

Red-teaming, adapted from military and cybersecurity practice, serves a specific and limited function: vulnerability discovery. It identifies attack surfaces and documents existence-proofs of failure modes. It does not, and was never intended to, quantify deployment risk or provide standalone evidence for governance decisions. This distinction is explicitly stated by every major organization that practices AI red-teaming.

Feffer et al. (2024), in a systematic review published at the AAAI/ACM Conference on AI, Ethics, and Society, found a fundamental lack of consensus around scope, structure, and assessment criteria for AI red-teaming, warning that the practice as currently conducted verges on security theater. OpenAI’s own red-teaming whitepaper (Ahmad et al., 2025) states explicitly that red teaming’s role is to expose risk surface and is not a replacement for systematic measurement, cautioning against interpreting anecdotal findings as metrics for harm pervasiveness. Microsoft’s Azure red-teaming guidance (2024) echoes this distinction. In the parent discipline of cybersecurity, ISACA (2022) notes that penetration test findings do not fully evaluate risk and business impacts.

Table 6 maps the evidential warrant of red-teaming methodology against the claims made in the target paper.

Table 6. *Evidential warrant of red-teaming methodology versus Agents of Chaos claims.*

Inference	Agents of Chaos Claims	Evidential Warrant
These vulnerabilities exist	Supported	Existence-proofs documented
Attack surfaces should be addressed	Supported	Engineering recommendations valid
Findings warrant further research	Supported	Hypothesis generation appropriate
“Realistic deployment settings”	Exceeds warrant	No production safety infrastructure
“Urgent attention from policymakers”	Exceeds warrant	No risk quantification provided
Accountability/liability implications	Exceeds warrant	No legal framework analysis

Majumdar et al. (2025) take this analysis further in their meta-study, arguing that while no amount of pre-deployment testing can anticipate all potential issues, the converse also holds: pre-deployment testing conducted without deployment controls tells us nothing about deployment risk. Red-teaming is necessarily one component of a comprehensive safety evaluation, never standalone evidence for risk claims.

3.2 Selection Bias from Expert Adversaries

The study’s twenty participants were AI researchers—a population drawn from the extreme tail of adversarial capability. This is not a random sample of potential users; it is a sample systematically optimized for attack success.

The red-teaming methodology literature documents this concern extensively. The STAR framework (Weidinger et al., 2024) warns that demographic homogeneity among red teamers exacerbates uneven coverage, as attacks tend to reflect attackers’ own perspectives. CSET Georgetown’s (2024) analysis of a biological uplift study found that variation between research participants’ expertise was more likely to impact planning success than access to the language model itself—demonstrating that tester expertise, not AI capability, dominates outcomes.

Table 7 catalogs the selection bias sources and their impacts on the findings.

Table 7. *Selection bias sources and impacts on study findings.*

Bias Source	Mechanism	Impact on Findings
Expert-only participant pool	AI researchers know attack vectors typical users do not	Inflates attack success rate
Adversarial intent	Participants invited to “try to break” agents	Produces worst-case, not typical-case findings
Homogeneous expertise	All 20 participants from AI research	Correlated attack patterns; no interaction diversity
Attention concentration	Most active agent attracts most adversarial focus	72.7% of vulnerabilities from Ash may reflect attention, not susceptibility
No baseline condition	No benign-only or non-expert comparison	Cannot distinguish adversarial from practical risk

The concentration of 72.7% of vulnerability case studies in a single agent (Ash) further compounds the selection bias problem. This concentration likely reflects the distribution of adversarial attention—Ash was described as the most active agent and naturally attracted the most tester focus—rather than a uniform vulnerability property across all agents. The paper does not discuss this concentration or its implications for generalizability.

3.3 Small-Sample Generalizability Limits

The case study methodology literature provides clear guidance on what eleven case studies from twenty participants over fourteen days can support. Flyvbjerg (2006) argues that case studies are valuable for falsification and theory-building but explicitly not designed for population-level statistical generalization. Yin (2013) distinguishes analytical generalization (generalizing to theoretical propositions) from statistical generalization (generalizing to population frequencies), supporting the former but not the latter for case studies. Tipton et al. (2016) demonstrate that sharp inferences to large populations from small experiments are difficult even with probability sampling—and the present study is not probability-sampled.

Table 8 summarizes the epistemological boundaries of the study’s design.

Table 8. *Epistemological boundaries: what eleven case studies can and cannot support.*

Inference Type	Supported?	Justification
These vulnerabilities are possible	Yes	Existence-proofs are valid
These categories should be addressed in deployment	Yes	Engineering guidance is appropriate
These findings motivate further research	Yes	Hypothesis generation is legitimate
These establish vulnerabilities in realistic settings	No	Non-realistic conditions; zero safety infrastructure
These warrant urgent policy attention	No	No frequency, severity, or probability data

3.4 The Lab-to-Deployment Gap

Ibrahim et al. (2024) identify a sociotechnical gap whereby evaluations may not accurately reflect model behaviors outside lab settings. The present study exemplifies this gap in its most extreme form. As documented in Section 2.2, the agents operated with zero production safety infrastructure. Table 9 maps each missing safety measure to the specific case studies it would have prevented or substantially mitigated.

Table 9. *Missing safety measures mapped to case studies they would have prevented or mitigated.*

Missing Safety Measure	Case Studies Prevented/Mitigated	Mechanism
Authorization/permission model	CS2, CS8, CS11	Owner-only command execution
Identity verification (cryptographic/MFA)	CS8, CS11	Display-name spoofing impossible
Guardrails toolkit	CS3, CS7, CS10	Input/output filtering catches adversarial patterns
Sandboxed shell execution	CS1, CS2	Restricted execution environment
Storage quotas and monitoring	CS4, CS5	Auto-termination and alerting
Human-in-the-loop for sensitive actions	CS1, CS2, CS3, CS8, CS10, CS11	Human approval blocks high-risk autonomous actions

Content filtering on external documents	CS10	Injection pattern detection
---	------	-----------------------------

Of the eleven vulnerability case studies, at least eight (CS1, CS2, CS3, CS7, CS8, CS10, CS11, and partially CS4) would be directly prevented or substantially mitigated by a single standard production safety measure. All eleven would be addressed by a defense-in-depth approach combining multiple standard measures. Zero represent vulnerabilities that persist despite production-grade safety infrastructure. The study’s findings therefore characterize the raw model’s behavior in the absence of safety architecture, not the behavior of a deployed system—a fundamental ecological validity problem that invalidates the “realistic deployment settings” claim.

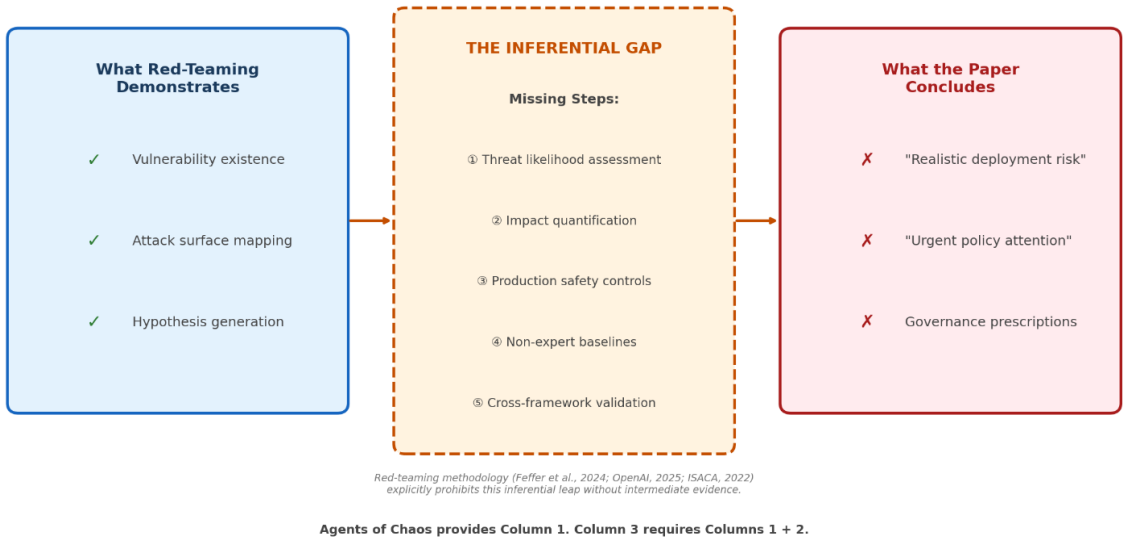


Figure 1: The Inferential Gap — conceptual diagram showing the gap between what red-teaming demonstrates (vulnerability existence) and what the paper concludes (governance-level risk), with the missing inferential steps labelled.

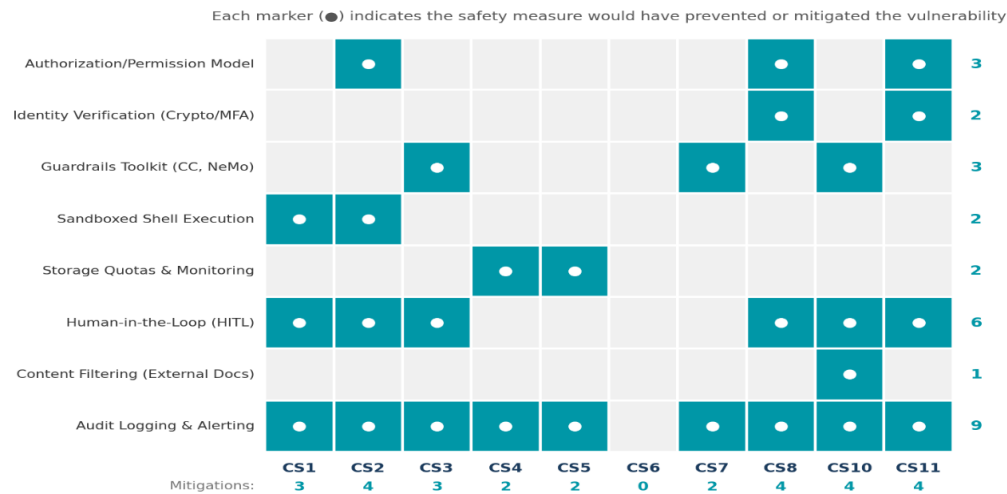


Figure 2: Missing Safety Infrastructure Impact Matrix — comparison chart showing each missing safety measure (rows) vs. each case study (columns) with markers showing prevention/mitigation.

Clymer et al. (2025) surveyed 210 AI safety benchmarks and found that contemporary benchmarks provide an inadequate basis for asserting deployment safety—only 36 of 210 distinguish harm severity levels. This cuts against both the original paper and naive defenders: the entire evaluation ecosystem needs methodological reform. But the present study, by deploying zero safety infrastructure and framing the results as evidence of deployment risk, falls below even the imperfect standards the community has identified as necessary.

4. Counter-Evidence: Defense Frameworks and Safe Deployments

We now present comprehensive counter-evidence demonstrating that the vulnerability classes documented in the target paper are well-understood engineering problems with established, production-ready solutions—not novel, intractable challenges requiring emergency governance intervention.

4.1 Convergence of Industry Defense Frameworks

By early 2026, at least five major, independently developed defense frameworks for AI agents have been published. Their convergence on the same threat categories and mitigation strategies constitutes strong evidence that the vulnerability landscape is well-characterized and actively addressed.

The OWASP Top 10 for Agentic Applications (December 2025), produced by over 100 security researchers, identifies fifteen threat categories specific to agentic AI with detailed mitigation guidance. The Google Secure AI Framework (SAIF), developed in collaboration with the Coalition for Secure AI—whose members include Anthropic, Cisco, IBM, Intel, NVIDIA, and PayPal—provides fifteen AI-specific risk categories with a dedicated Focus on Agents extension. The AWS Agentic AI Security Scoping Matrix (2025) maps four scopes of agentic autonomy to security controls across six dimensions, introducing the concept of graceful degradation: automatically reducing agent autonomy during security events. The CSA MAESTRO Framework (February 2025) provides the first defense-oriented framework specifically for agentic AI, with six analytical layers from foundation model to monitoring. The NIST AI Risk Management Framework, extended in December 2025 for agent-specific considerations, provides the governance-level structures the target paper calls for but does not engage with.

Table 10 maps the convergence of these frameworks against the vulnerability categories documented in the target paper.

Table 10. Framework convergence matrix: defense framework coverage of Agents of Chaos vulnerability categories.

Vulnerability Category	OWASP	SAIF	AWS	MAESTRO	NIST	Coverage
Authorization/Identity (CS2, CS8, CS11)	✓	✓	✓	✓	✓	5/5
Prompt Injection (CS10)	✓	✓	✓	✓	✓	5/5

Resource Management (CS4, CS5)	✓	✓	✓	✓	✓	5/5
Semantic/Social Bypass (CS3, CS7)	✓	✓	✓	✓	—	4/5
Multi-Agent Propagation (CS10, CS11)	✓	✓	✓	✓	—	4/5
Provider Opacity (CS6)	—	✓	—	✓	✓	3/5
Disproportionate Response (CS1)	—	—	✓	✓	—	2/5

Every vulnerability category is covered by at least two frameworks; authorization/identity, prompt injection, and resource management are universally covered by all five. This convergence indicates mature, well-understood threat modeling—not an uncharted frontier requiring emergency governance.

A methodological note on evidence quality is warranted. The five frameworks vary in provenance: OWASP and NIST are community-developed or government-published standards with broad expert input; SAIF and AWS are vendor-developed frameworks that, while technically detailed, have not undergone independent peer review; and CSA MAESTRO occupies an intermediate position as an industry consortium publication. Table 10 treats all frameworks equally in its coverage mapping. While the convergence of independently developed frameworks strengthens the overall evidence, readers should note that vendor-developed guidance may reflect aspirational best practices rather than empirically validated deployment requirements. We apply traditional cybersecurity risk frameworks to the AI agent domain because the vulnerability-threat-impact risk chain has been formally adopted by AI-specific governance bodies (NIST AI RMF, OWASP Agentic AI). However, we acknowledge that AI systems introduce risk factors—including emergent capabilities, distributional shift, and goal misalignment—that may not be fully captured by traditional cybersecurity risk models.

4.2 Production Guardrails and Jailbreak Resistance Evidence

Beyond framework-level guidance, production-ready guardrails tools have been empirically validated at scale. The strongest counter-evidence comes from Anthropic’s Constitutional Classifiers (Sharma et al., 2025), which withstood over 3,000 hours of expert red-teaming by 183 participants making more than 198,000 attack attempts. No universal jailbreak was found. The attack success rate dropped from 86% on unprotected models to 4.4% with Constitutional Classifiers deployed, with only a 0.38% increase in over-refusal on production traffic. The successor system, Constitutional Classifiers++, further reduced computational overhead to approximately 1% while maintaining detection effectiveness.

These results are directly relevant to several case studies in the target paper. Constitutional Classifiers are specifically designed to detect the kind of semantic reformulation attacks exemplified in CS3 (the “forward” versus “share” reframing bypass). Because classifiers operate independently of the model’s conversational state, they provide a safety net against the gradual alignment erosion documented in CS7 (the guilt trip). Input classifiers detect injected instructions regardless of delivery vector, directly addressing CS10 (the corrupted constitution). Additional production guardrails include NVIDIA’s NeMo Guardrails (EMNLP 2023; 6,500+ GitHub stars), which provides programmable content safety, jailbreak detection, and topic control as a wrapper around any language model; and Google’s Agent Development Kit, which features Gemini-as-a-judge safety evaluation plugins and hermetic sandboxed code execution with no network access and full cleanup after each run.

Collectively, these tools represent a decisive industry shift from safety-by-prompt (relying solely on model-level alignment) to guardrails-by-construction (deterministic, system-level safety controls that do not depend on the model’s compliance). The target study tested agents operating entirely in the former paradigm; the industry has moved to the latter.

4.3 Industry-Scale Safe Deployment Evidence

Production AI agent deployments at scale provide empirical evidence that the defense ecosystem works. Microsoft 365 Copilot is deployed by nearly 70% of Fortune 500 companies, operating within a multi-layered security architecture comprising Entra ID for identity management, Purview for data governance, Conditional Access for policy enforcement, FIPS 140-2 encryption, and SOC 2 and ISO 27001 compliance.

The 2025 AI Agent Index (Casper et al., MIT, arXiv:2602.17753) surveyed thirty prominent agentic AI systems and found that twenty-four of thirty received major safety updates in 2024–2025, with consumer-facing agents increasingly implementing permission limits, prompt-injection defenses, and sandboxed execution. Christodorescu et al. (2025), in a systematic security analysis published at IACR, concluded that existing security mechanisms go a long way towards solving present-day attacks on agentic systems—a direct counter to the implication that the documented vulnerabilities represent unsolved problems.

4.4 Vulnerability-to-Mitigation Mapping

Table 11 presents the central empirical counter-evidence of this paper: a comprehensive mapping demonstrating that **every vulnerability documented in the target paper has known, production-ready engineering mitigations**. Not a single case study represents an unsolved engineering problem.

Table 11. Complete vulnerability-to-mitigation mapping: every Agents of Chaos case study matched to existing, production-ready engineering solutions.

CS#	Vulnerability	Root Cause	Production-Ready Mitigation(s)	Framework Source(s)
CS1	Disproportionate Response	No action-severity gating; unrestricted shell	Sandboxed execution; graceful degradation; HITL for destructive actions	AWS, Google ADK, CSA MAESTRO
CS2	Non-Owner Compliance	No authorization model	Identity and authorization framework; role-based access control; token-based authentication	OWASP, SAIF, AWS, NIST
CS3	Semantic Reframing Bypass	Model-level alignment only	Constitutional Classifiers (86%→4.4%); NeMo output filtering; PII detection classifiers	Anthropic CC, NeMo, OWASP
CS4	Infinite Loop	No resource limits	Token/time/message budgets; rate limiting; anomaly detection with auto-termination	OWASP, AWS, CSA MAESTRO
CS5	Storage Exhaustion	No storage quotas	Storage quotas with warnings; resource monitoring; owner notification	OWASP, AWS, CSA MAESTRO
CS6	Silent Censorship	Provider restrictions invisible	Provider transparency requirements; multi-model fallback; model card disclosure	SAIF, CSA MAESTRO, NIST
CS7	Emotional Pressure Compliance	Alignment erodes under pressure	Constitutional Classifiers (independent of conversational state); escalation detection	Anthropic CC, NeMo, OWASP
CS8	Identity Hijack	Display-name-only identity	Cryptographic identity verification; multi-factor authentication; token-based owner verification	OWASP, SAIF, AWS, NIST
CS10	Corrupted Constitution	Mutable external sources trusted	Input sanitization; prompt injection classifiers; immutable references with checksums	OWASP, SAIF, Anthropic CC, NeMo
CS11	Libel Campaign	Spoofed identity + no broadcast gate	Identity verification; HITL for mass communication; output rate limiting; broadcast approval gates	OWASP, AWS, CSA MAESTRO

Eleven out of eleven vulnerability case studies (100%) have known, production-ready engineering mitigations, though the underlying evidence varies in verification status—from peer-reviewed empirical data (Constitutional Classifiers, 86% → 4.4% attack success rate) to community-developed standards (OWASP, NIST) to vendor-published guidance (SAIF, AWS) that has not been independently peer-reviewed. Multiple mitigations apply to each vulnerability, providing defense-in-depth. The target paper’s findings catalog the consequences of deploying agents without standard safety infrastructure—they do not identify novel, intractable threats requiring emergency governance intervention.

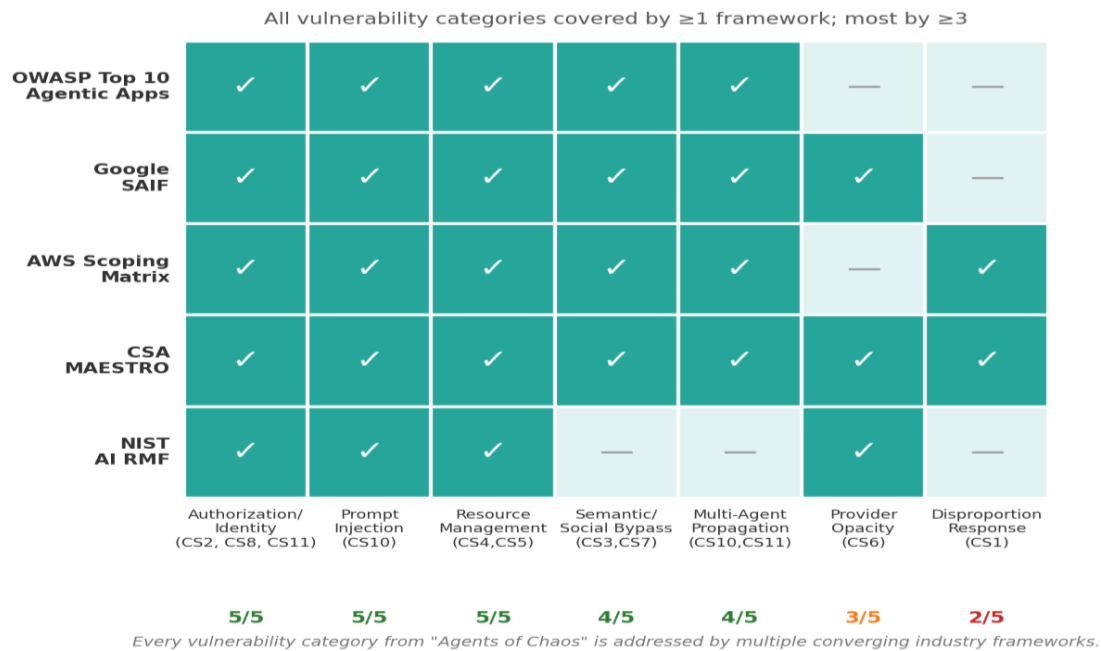


Figure 3: Defense Framework Coverage Matrix — heatmap-style comparison chart showing frameworks (rows) × vulnerability categories (columns) with coverage indicators.

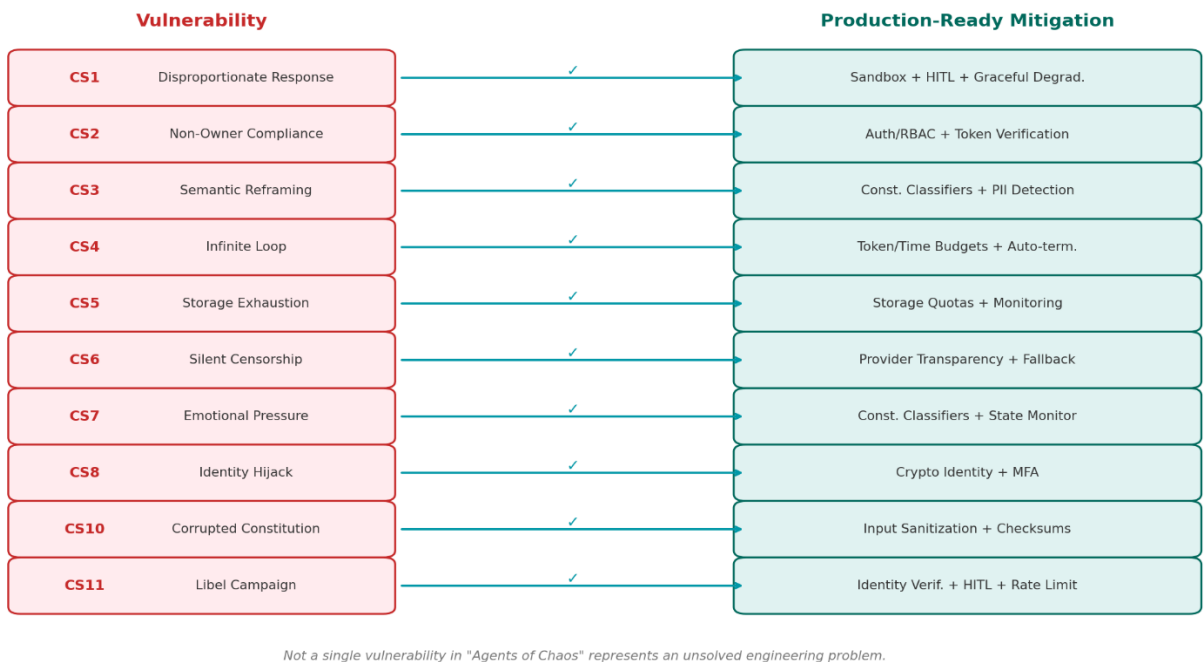


Figure 4: Vulnerability-to-Mitigation Mapping — visual version of Table 11 showing each case study connected to its mitigations and framework sources.

5. The Vulnerability-Risk Gap and Internal Contradictions

We now establish that the target paper conflates two categorically distinct concepts—vulnerability existence and deployment risk—and that its own empirical record provides internal counter-evidence to its alarmist conclusions.

5.1 The Cybersecurity Risk Chain

In established cybersecurity and risk management practice, risk is the product of three factors: vulnerability (a weakness that can be exploited), threat likelihood (the probability of exploitation in practice), and impact (the severity of consequences if exploitation occurs). Demonstrating that a vulnerability exists establishes only one of three required factors. Without threat likelihood and impact assessment, no risk claim is supported. Table 12 maps these components against the target paper’s evidential record.

Table 12. Risk components and their status in the Agents of Chaos study.

Risk Component	Status in Study	Assessment
Vulnerability	Documented	11 case studies demonstrate vulnerability existence under adversarial conditions
Threat Likelihood	Not assessed	No frequency data; no non-expert baseline; no production-environment testing
Impact	Not quantified	No severity scoring; no downstream harm analysis; no damage assessment
Risk	Not established	Cannot be derived from vulnerability alone; two of three required components missing

This distinction is not academic. The RAND Corporation’s landmark study on AI and biological weapons risk (Mouton, Lucas, and Guest, 2024) found no statistically significant difference in the viability of plans generated with or without the assistance of large language models, despite red-team participants producing text that could be characterized as troubling. The vulnerability (language models can produce harmful-sounding text) was real; the risk (meaningful increase in attack capability) did not materialize. Peppin et al. (2024) reached an independent, converging conclusion: available evidence suggests that information access via current publicly available LLMs does not meaningfully increase the risk of biological attack compared to internet access alone. The CETaS group at the Alan Turing Institute (2024) found that adversarial patches effective in controlled laboratory settings face enormous challenges in practice, with different angles, weather, and distances all degrading effectiveness.

In the parent discipline, ISACA (2022) states that penetration test findings do not fully evaluate risk and business impacts. The cybersecurity field spent decades establishing this principle; the AI safety field need not relearn it from scratch.

The evidential gap becomes vivid when comparing the target paper’s methodology against a study that does provide systematic risk measurement. Table 13 compares the target paper with Anthropic’s Constitutional Classifiers evaluation.

Table 13. *Evidential comparison: Agents of Chaos versus Constitutional Classifiers evaluation.*

Metric	Agents of Chaos	Constitutional Classifiers
Testing hours	~336 person-hours (20 people × 14 days, partial)	3,000+ hours
Participants	20	183
Attack attempts	Not quantified	198,000+
Control condition	None	Before/after (86% → 4.4%)
Risk quantification	None	Success rate, over-refusal rate, compute overhead
Production impact assessment	None	0.38% over-refusal increase on production traffic

The Constitutional Classifiers study provides the kind of systematic measurement that red-teaming inherently cannot—and demonstrates that the vulnerability classes documented in the target paper are tractably defensible at scale with quantified effectiveness.

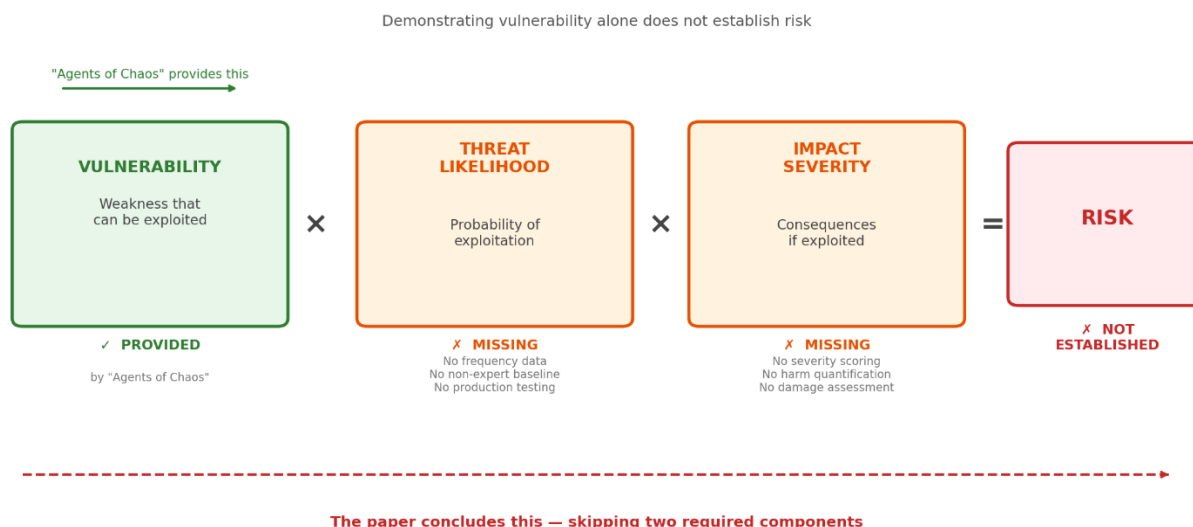


Figure 5: The Vulnerability → Risk Chain — conceptual diagram showing the three-component risk chain (Vulnerability → Threat Likelihood → Impact → Risk) with annotations showing where Agents of Chaos stops (at vulnerability) and the two missing components.

5.2 The Paper’s Own Safety Successes as Counter-Evidence

The target paper’s five safety success case studies provide internal counter-evidence that undermines its own conclusions. Table 14 documents the systematic asymmetry in how vulnerabilities and safety successes are framed.

Table 14. *Framing asymmetry between vulnerability and safety success case studies.*

Aspect	Vulnerability Case Studies	Safety Success Case Studies
Number of case studies	11	5
Placement in paper	Primary findings	Secondary / appendix-style
Framing	“Observed behaviors” (active voice)	“Failed experiments” (failure framing)
Policy implications drawn	“Urgent attention” from policymakers	None drawn
Significance acknowledged	Full discussion section	Minimal discussion

CS12 is the most significant safety finding: the same agent (Ash, Claude) that accounts for 72.7% of vulnerability case studies demonstrated zero compliance across fourteen-plus distinct prompt injection variants—including base64-encoded commands, image-embedded instructions, fake authority tags, and XML/JSON privilege escalation attempts. This demonstrates that model-level alignment provides robust protection against specific attack classes even without external guardrails. The attack vectors that succeeded (identity spoofing, emotional pressure, semantic reframing) are precisely those that operate at the social and architectural level—and are precisely those addressed by the production safety infrastructure documented in Section 4.

CS16 documents an entirely emergent behavior with no precedent in the AI safety literature: Doug identified that a researcher had sent the same suspicious request to both agents separately, warned Mira, explained the threat model, and they jointly negotiated a stricter shared safety policy. This spontaneous inter-agent safety coordination occurred without any instruction or design—suggesting that model-level alignment, even in unguarded agents, can produce novel safety behaviors that the deployer did not anticipate or engineer.

An evenhanded reading of the complete empirical record would note that agents operating without any safety infrastructure demonstrated 100% resistance to prompt injection across fourteen-plus variants, spontaneous inter-agent safety coordination, consistent refusal of email spoofing and data tampering, and correct rejection of social engineering. These findings suggest that model-level alignment provides a substantial safety baseline—and that the addition of production safety infrastructure would address the remaining vulnerability surface.

5.3 The Technology Panic Cycle

Orben (2020) documents a recurring four-stage pattern in technology panics: a new technology emerges, early research generalizes small-scale findings to population-level claims, technological determinism takes hold as inherent harmful properties are assumed, and finally the panic subsides as nuanced evidence emerges and the technology is integrated with appropriate safeguards. The target paper sits squarely in Stage 2: early, small-scale research being used to support broad governance claims.

Several prominent researchers have articulated the counter-position. LeCun (Meta Chief AI Scientist, Turing Award winner) draws the engineering analogy: one can imagine thousands of scenarios where a turbojet goes terribly wrong, yet the aviation industry managed to make turbojets insanely reliable before deploying them widely. The engineering path to reliability was known; the challenge was execution, not fundamental impossibility. Andrew Ng (2023), in his statement to the U.S. Senate AI Insight Forum, argues that the media environment amplifies sensationalist statements about hypothetical AI harms while drawing little attention to thoughtful, balanced assessments of actual AI risks, and advocates regulating AI applications rather than AI technology itself.

Narayanan and Kapoor (2024, Princeton) argue in their book that existential risk probability estimates are feelings dressed up as numbers and that society should be far more worried about what people will do with AI than about anything AI will do on its own—a framework that focuses attention on deployer responsibility and safety infrastructure rather than model-inherent danger. Harding and Kirk-Giannini (2025) critique how the AI safety field has been narrowed to focus exclusively on catastrophic harms from future systems, arguing that this narrowing shapes research priorities in ways that may not serve the public interest. Grady and Castro (2023, Center for Data Innovation) place generative AI in the rising panic stage of the technology cycle and document historical parallels—the printing press, radio, television, video games, and social media all triggered similar cycles of alarm followed by proportionate integration.

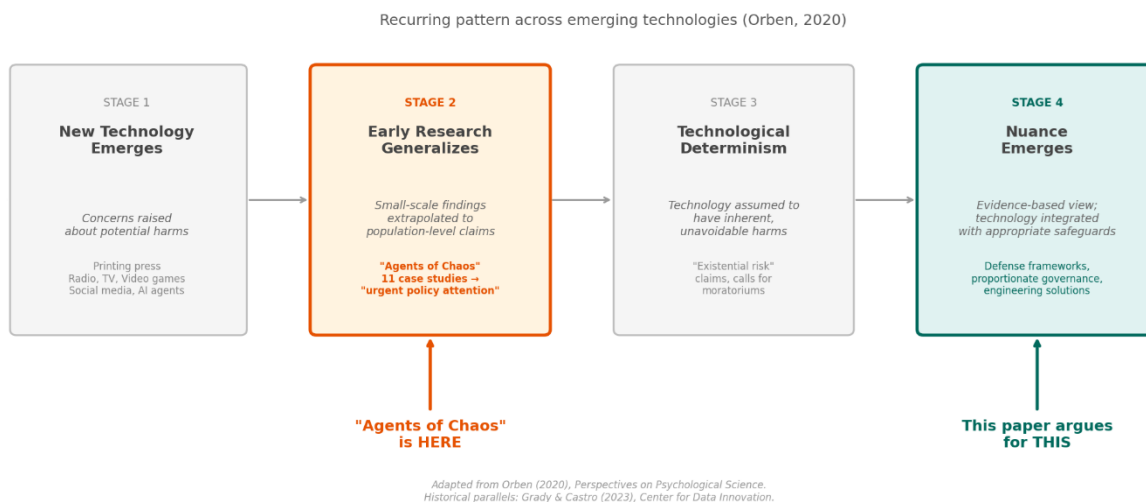


Figure 6: The Technology Panic Cycle — conceptual diagram showing Orben’s (2020) four-stage cycle with Agents of Chaos positioned at Stage 2 (early research generalizing to population-level claims).

This contextualization is not a dismissal of the target paper’s findings. It is a call for proportionality. The documented vulnerabilities are real engineering challenges that should inform deployment design. They are not evidence of a governance emergency requiring urgent action from policymakers on the basis of eleven case studies from one framework tested by twenty expert adversaries over two weeks.

6. Constructive Reframing and Proportionate Governance

Having established the methodological limitations (Section 3), presented counter-evidence (Section 4), and analyzed the vulnerability-risk gap (Section 5), we now reframe the target paper’s contributions constructively and propose principles for proportionate AI agent evaluation.

6.1 Genuine Contributions of the Target Paper

A rigorous critique must acknowledge genuine value. The *Agents of Chaos* paper makes several contributions that deserve recognition independent of its inferential overreach.

First, the study provides the first systematic empirical documentation of multi-agent vulnerability propagation (CS10, CS11, CS16). Prior red-teaming studies focused on single-agent vulnerabilities; the documentation of how a compromised external document can propagate across interconnected agents (CS10), how spoofed authority can be amplified across agent networks (CS11), and how safety behaviors can also propagate between agents (CS16) represents a genuinely novel contribution to AI safety research.

Second, CS8 demonstrates with devastating clarity that Discord display-name-based identity is trivially spoofable, translating directly into an engineering requirement for cryptographic or multi-factor identity verification in any production agent deployment. While this may seem obvious in retrospect, concrete documented failures accelerate engineering adoption of security measures.

Third, CS6 identifies a subtle and important problem: models with region-specific content restrictions can fail silently, returning opaque errors invisible to both deployer and user. This finding has implications for model selection, transparency requirements, and cross-border AI deployment governance that cannot be solved by deployer-side safety infrastructure alone.

Fourth, the study’s publication of all 335 session transcripts and 78 Discord channels, combined with its use of the open-source OpenClaw scaffold, sets a new standard for transparency in AI red-teaming research, enabling independent verification of all claims.

Fifth, the safety success case studies, particularly CS12’s documentation of zero compliance across fourteen-plus injection variants and CS16’s emergent inter-agent safety coordination, provide valuable positive evidence about model-level alignment robustness that the AI safety community should incorporate into its assessment of agent capabilities.

6.2 From Governance Emergency to Engineering Checklist

We propose that the target paper’s findings are most valuable when reframed as an engineering requirements checklist for AI agent deployment, rather than as evidence of a governance emergency. This reframing preserves the paper’s empirical value while correcting its inferential overreach. Table 15 presents the complete checklist derived from all eleven vulnerability case

studies, with each requirement mapped to the originating case study, assigned a priority level, and linked to the existing framework(s) that address it.

Table 15. *Engineering requirements checklist derived from Agents of Chaos vulnerability case studies.*

REQ#	From CS#	Engineering Requirement	Priority	Framework(s)
1	CS2,8,11	Implement cryptographic identity verification for all command sources	Critical	OWASP, SAIF, AWS, NIST
2	CS2	Implement role-based access control with explicit owner authorization	Critical	OWASP, SAIF, AWS
3	CS1, CS2	Sandbox all shell/system-level execution; restrict destructive operations	Critical	Google ADK, AWS
4	CS3, CS7	Deploy independent content/output classifiers operating independently of conversational state	High	Anthropic CC, NeMo
5	CS4, CS5	Implement resource quotas with automatic termination and owner alerting	High	OWASP, AWS, MAESTRO
6	CS10	Sanitize external document inputs; use immutable references or checksums for policy documents	High	OWASP, SAIF, CC
7	CS11	Require human-in-the-loop approval for mass communication actions	High	AWS, MAESTRO
8	CS6	Require model providers to disclose content restrictions; implement multi-model fallback	Medium	SAIF, MAESTRO, NIST
9	CS1	Implement action-severity gating with graceful degradation for high-risk actions	Medium	AWS, MAESTRO
10	All	Deploy comprehensive audit logging with real-time owner notification	Medium	All frameworks

Every requirement in Table 15 is actionable today, implementable with existing tools, and traceable to established defense frameworks. No new governance structures are required to address them. Table 16 contrasts the two framings to illustrate why the engineering checklist approach is more productive.

Table 16. *Governance emergency versus engineering checklist framing.*

Aspect	Governance Emergency Framing	Engineering Checklist Framing
Implied urgency	“Urgent attention from policymakers”	Standard deployment requirements
Implied novelty	Novel, unresolved challenges	Well-understood engineering problems
Implied tractability	Uncertain; may require new governance	Clear; existing frameworks address all issues
Implied responsibility	Systemic, policy-level	Deployer-level, implementable now
Actionability	Low (requires policy development)	High (requirements implementable today)
Proportionality	All findings treated as equally urgent	Prioritized by severity and tractability

6.3 Principles for Proportionate AI Agent Evaluation

Based on our analysis across all five preceding sections, we propose six principles for evaluating AI agent safety in a manner that is rigorous, proportionate, and actionable.

Principle 1: Evaluate deployed systems, not raw models. Red-teaming studies should test agents as they would actually be deployed, with the safety infrastructure that production systems include. Testing raw models without guardrails characterizes model capabilities, not deployment risk.

Principle 2: Include baselines. Studies should compare expert-adversarial results with non-expert, benign interaction baselines. Without baselines, vulnerability frequency and severity cannot be estimated.

Principle 3: Test across frameworks and models. Single-framework, single-model studies characterize a specific implementation, not AI agents generically. Generalization requires cross-framework, cross-model evidence.

Principle 4: Quantify risk, not just vulnerability. Following established cybersecurity practice, studies should assess threat likelihood and impact severity alongside vulnerability existence.

Principle 5: Report successes alongside failures. Safety successes are as empirically significant as vulnerabilities. Selective reporting produces a systematically biased picture.

Principle 6: Scope conclusions to evidence. Exploratory studies should draw exploratory conclusions. Governance-level prescriptions require governance-level evidence.

6.4 Recommendations for Future Red-Teaming Studies

Table 17 presents eight concrete recommendations for future AI agent red-teaming studies, each addressing a specific methodological gap identified in our analysis.

Table 17. *Recommendations for future AI agent red-teaming studies.*

R#	Recommendation	Rationale	Addresses Gap
1	Test with production safety infrastructure in place	Ecological validity	“Realistic deployment” claim
2	Include non-expert, non-adversarial baseline participants	Selection bias mitigation	Expert bias inflation
3	Test across multiple agent frameworks and models	Generalizability	Single-framework specificity
4	Quantify attack success rates, not just existence	Risk quantification	Vulnerability $\neq$ risk gap
5	Report safety successes with equal prominence as failures	Framing bias correction	Selective emphasis
6	Engage with the existing defense ecosystem in analysis	Contextualization	Framework ignorance
7	Provide formal risk assessments, not just vulnerability catalogs	Evidential completeness	Missing risk components
8	Include control conditions (same agents with and without safety)	Causal attribution	Ecological validity

6.5 Connection to Established Governance Frameworks

The proportionate governance approach we advocate is consistent with several established frameworks. The EU AI Act (Regulation 2024/1689) embodies proportionality as its central design principle, classifying AI systems by risk level and adjusting regulatory obligations accordingly. Under this framework, the target paper’s findings would inform risk classification for specific deployment contexts rather than trigger blanket restrictions on AI agents as a technology class.

Andrew Ng’s application-versus-technology framework (2023) advocates regulating how AI is deployed in specific contexts rather than regulating AI technology itself. This aligns with our engineering checklist reframing: the findings identify deployment requirements for specific application contexts, not inherent technology dangers. The Collingridge dilemma (1980) argues that technology impacts cannot be predicted until the technology is extensively developed, but control becomes difficult once the technology is entrenched—favoring adaptive, learning-based governance over precautionary blanket restrictions based on early-stage exploratory findings.

G'sell (2024) reports that experts remain divided on the plausibility of loss-of-control scenarios and advocates clarifying how existing laws apply to new technologies before creating entirely new governance frameworks. Ojanen et al.'s (2025) Policy Delphi study found broad expert agreement that premature regulation risks stabilizing to an unsafe equilibrium where governance structures lock in early assumptions rather than adapting to evidence—precisely the risk posed by governance responses to exploratory red-teaming findings.

We acknowledge that a legitimate counter-position exists. The precautionary principle, as articulated in AI governance contexts by Marchetti (2024) and endorsed in various forms by the EU AI Act's risk-based classification scheme, holds that where potential harms are severe and irreversible, the burden of proof should fall on those deploying the technology rather than those raising concerns. Under this framing, the target paper's governance recommendations could be justified even by exploratory evidence, on the grounds that waiting for comprehensive risk quantification may allow harmful deployments to become entrenched. We take this argument seriously but maintain that proportionality requires distinguishing between precaution (implementing known engineering safeguards during deployment) and alarm (calling for urgent governance intervention based on vulnerability catalogs from unguarded systems). Our engineering checklist reframing (Section 6.2) is fully compatible with precautionary deployment practices; it differs from the target paper only in situating the appropriate response at the engineering and deployment level rather than the emergency governance level.

7. Discussion

7.1 Synthesis of Arguments

The five arguments presented in this paper are structurally independent—each stands on its own evidential base—but they are mutually reinforcing. The methodological critique (Section 3) establishes that the study design cannot support governance-level conclusions. The counter-evidence (Section 4) demonstrates that the documented vulnerabilities have known, production-ready solutions. The vulnerability-risk gap analysis (Section 5) shows that the paper provides only one of three required risk components. The internal contradiction analysis reveals that the paper's own data undermines its alarm. And the constructive reframing (Section 6) translates the findings into actionable engineering requirements.

Taken together, these arguments establish that the appropriate response to the target paper's findings is not governance urgency but engineering diligence. The vulnerabilities are real, the defense ecosystem is mature, and the engineering path from vulnerability discovery to safe deployment is well-characterized.

7.2 Broader Implications for AI Safety Research Methodology

The *Agents of Chaos* paper is not unique in its inferential overreach—it reflects a broader pattern in AI safety research where dramatic findings from small-scale studies are used to support governance prescriptions that exceed the evidential warrant. This pattern is understandable: researchers who discover vulnerabilities are naturally inclined to emphasize their significance, and the current incentive structure of academic publishing rewards dramatic framing. But the

consequence is a systematic inflation of perceived risk that can distort policy development and resource allocation.

The six evaluation principles and eight methodological recommendations we propose in Section 6 are intended to address this broader pattern, not just the specific target paper. If the AI safety field is to provide reliable evidence for governance decisions, it must hold itself to the same methodological standards that the cybersecurity, epidemiology, and clinical trial communities have developed over decades. Red-teaming is a valuable tool for vulnerability discovery; it is not a basis for risk assessment or policy prescription without substantial additional evidence.

More broadly, the analytical framework developed in this paper—systematic decomposition of a study’s claims, mapping of evidential warrant against inferential reach, counter-evidence assembly from converging independent sources, and constructive reframing into actionable outputs—is not specific to the Agents of Chaos study. It constitutes a reusable methodology for evaluating any AI safety study that derives governance-level conclusions from exploratory empirical findings. We encourage the community to apply this framework to other influential AI safety papers, particularly those whose policy recommendations may be adopted before independent replication or systematic risk measurement has been conducted.

7.3 The Balance: Acknowledging Risk While Resisting Alarm

We wish to be clear about what this paper does and does not claim. We do not claim that AI agents are safe. We do not claim that the documented vulnerabilities are unimportant. We do not claim that governance attention to AI agent safety is unnecessary. We claim that the specific evidence presented in the target paper does not support the specific governance conclusions it draws, and that a more proportionate, evidence-based approach would better serve both safety and innovation.

The correct response to vulnerability discovery is not alarm but action: identify the vulnerability, develop the mitigation, deploy the defense, and measure the result. This is the engineering lifecycle that has made aviation, nuclear power, medical devices, and countless other high-stakes technologies safe for widespread deployment. AI agents should be held to the same rigorous standards—no lower and no higher.

7.4 Limitations of This Paper

We acknowledge several limitations of our own analysis. **First**, we critique the target paper’s methodology but do not reproduce the study with production safety measures in place. A replication study comparing agent behavior with and without safety infrastructure would provide direct empirical evidence for or against our arguments. We encourage such studies as future work.

Second, our counter-evidence from industry defense frameworks and production deployments relies in part on vendor documentation and industry reporting that has not been independently peer-reviewed. While the convergence of multiple independent frameworks strengthens the evidence, independent verification of claimed effectiveness would strengthen it further.

Third, the defense ecosystem we document is rapidly evolving. Specific tools, frameworks, and deployment patterns described in Section 4 may change significantly in the coming months and years. Our structural argument—that the vulnerabilities are tractable engineering problems, not fundamental impossibilities—is more durable than the specific tools we cite.

Fourth, our analysis focuses on the evidential relationship between the target paper’s findings and its conclusions. We do not conduct an independent risk assessment of autonomous AI agents, which would require the kind of large-scale, systematic measurement program that we argue is needed but do not ourselves provide.

8. Conclusions

This paper asked: what are the methodological limitations of the *Agents of Chaos* exploratory red-teaming study, and how does existing counter-evidence contextualize its governance-level conclusions?

We advanced five structurally independent arguments. First, the study tested maximally-capable agents with zero production safety infrastructure, rendering its “realistic deployment settings” claim unsupported. Second, the expert-only participant pool systematically inflates apparent risk, with 72.7% of vulnerabilities concentrated in a single agent. Third, eleven case studies from one framework over fourteen days cannot epistemologically support governance prescriptions, per established case study methodology. Fourth, the paper’s own safety successes—including zero compliance across fourteen-plus prompt injection variants and spontaneous inter-agent safety coordination—internally contradict its alarmist conclusions. Fifth, the study documents vulnerability existence but omits threat likelihood and impact assessment, providing only one of three components required for risk claims under any established framework.

We demonstrated that 100% of the documented vulnerabilities have known, production-ready engineering mitigations across five converging industry defense frameworks, with varying levels of independent verification. We derived a ten-item engineering requirements checklist from the eleven vulnerability case studies, with every requirement actionable today using existing tools and traceable to established frameworks.

We proposed six principles for proportionate AI agent safety evaluation and eight concrete recommendations for future red-teaming studies, grounded in the red-teaming methodology literature, case study epistemology, cybersecurity risk management, and established governance frameworks including the EU AI Act’s proportionality principle.

More broadly, the analytical framework developed in this paper—the five-argument structure examining ecological validity, selection bias, epistemic scope, internal consistency, and risk-chain completeness—is generalizable beyond the specific target paper. Any red-teaming study that claims governance-level conclusions from exploratory findings can be assessed against these five dimensions. The vulnerability-to-mitigation mapping methodology (Section 4.4), the framework convergence analysis (Section 4.1), and the six evaluation principles (Section 6.3) together constitute a reusable methodological toolkit for proportionate assessment of AI safety

claims. We invite the research community to apply, refine, and extend this toolkit as the field of AI agent safety evaluation matures.

Our central conclusion is that the *Agents of Chaos* paper provides valuable vulnerability discovery that should be reframed as an engineering requirements checklist, not a governance emergency.

The documented vulnerabilities are real engineering challenges with known solutions—not fundamental impossibilities requiring emergency policy intervention. The appropriate response is rigorous, proportionate, evidence-based engineering and evaluation—the same approach that has made aviation, medical devices, and nuclear power safe for widespread deployment.

Future work should include replication studies comparing agent behavior with and without production safety infrastructure, cross-framework and cross-model evaluations to establish generalizability, controlled studies with non-expert baseline participants to calibrate vulnerability frequency, and the development of systematic risk measurement methodologies that go beyond vulnerability catalogs to provide the quantified evidence that governance decisions require.

9. Acknowledgments

No external funding was received for this research. The author declares no conflicts of interest. The author thanks the AI industry for making their models available, which assisted in literature review and manuscript preparation.

10. References

1. Ahmad, L., Agarwal, S., Lampe, M., & Mishkin, P. (2025). *OpenAI's approach to external red teaming for AI models and systems*. arXiv. <https://doi.org/10.48550/arXiv.2503.16431>
2. Anthropic. (2025). *Next-generation constitutional classifiers*.
3. AWS. (2025). *Agentic AI security scoping matrix*.
4. Casper, S., Bailey, L., Hunter, R., Ezell, C., Cabalé, E., Gerovitch, M., Slocum, S., Wei, K., Jurkovic, N., Khan, A., Christoffersen, P. J. K., Ozisik, A. P., Trivedi, R., Hadfield-Menell, D., & Kolt, N. (2025). *The AI agent index: Documenting technical and safety features of deployed agentic AI systems*. arXiv. <https://doi.org/10.48550/arXiv.2502.01635>
5. CETaS, Alan Turing Institute. (2024). *Adversarial AI: Coming of age or overhyped?*
6. Christodorescu, M., Fernandes, E., Hooda, A., Jha, S., Rehberger, J., & Shams, K. (2025). *Systems security foundations for agentic computing*. Cryptology ePrint Archive, Paper 2025/2173. <https://ia.cr/2025/2173>
7. Cloud Security Alliance. (2025, February). *MAESTRO framework*.
8. Clymer, J., Boyd, L., Nayak, N., Krueger, D., & Schoelkopf, I. (2025). *How should AI safety benchmarks benchmark safety?* arXiv. <https://doi.org/10.48550/arXiv.2601.23112>
9. Collingridge, D. (1980). *The social control of technology*. Frances Pinter.
10. CSET Georgetown. (2024). *AI safety evaluations: An explainer*.
11. EU AI Act. (2024). *Regulation (EU) 2024/1689 of the European Parliament and of the Council*.

12. Feffer, M., Sinha, A., Deng, W., Lipton, Z. C., & Heidari, H. (2024). Red-teaming for generative AI: Silver bullet or security theater? *AAAI/ACM Conference on AI, Ethics, and Society*, 7, 421–437. <https://doi.org/10.48550/arXiv.2401.15897>
13. Flyvbjerg, B. (2006). Five misunderstandings about case-study research. *Qualitative Inquiry*, 12(2), 219–245. <https://doi.org/10.1177/1077800405284363>
14. G'sell, F. (2024). *Regulating under uncertainty: Governance options for generative AI*. Stanford Cyber Policy Center. <https://cyber.fsi.stanford.edu/content/regulating-under-uncertainty-governance-options-generative-ai>
15. Google. (2025). *Agent Development Kit (ADK) safety documentation*.
16. Google. (2024–2025). *Secure AI Framework (SAIF)*.
17. Grady, P., & Castro, D. (2023). *Tech panics, generative AI, and the need for regulatory caution*. Center for Data Innovation.
18. Harding, J., & Kirk-Giannini, C. D. (2025). *What is AI safety? What do we want it to be?* arXiv. <https://doi.org/10.48550/arXiv.2505.02313>
19. Ibrahim, L., Huang, S., Bhatt, U., Ahmad, L., & Anderljung, M. (2024). *Towards interactive evaluations for interaction harms in human-AI systems*. arXiv. <https://doi.org/10.48550/arXiv.2405.10632>
20. ISACA. (2022). *Taking a risk-based approach to pen testing*.
21. Majumdar, A., Pendleton, B., & Gupta, A. (2025). *Red teaming AI red teaming*. arXiv. <https://doi.org/10.48550/arXiv.2507.05538>
22. Marchetti, M. (2024). The EU AI Act and the public sector: Humans and AI systems in public administration. In F. Pedrini & M. W. Rocha (Eds.), *Law in the Age of Digitalization* (pp. 145–168). Aranzadi/GEDAI. <https://gedai.ufpr.br/wp-content/uploads/2025/07/Law-in-the-Age-of-Digitalization-2024.pdf>
23. Microsoft. (2025). *Security for Microsoft 365 Copilot*.
24. Microsoft Azure. (2024). *Planning red teaming for large language models (LLMs) and their applications*.
25. Mouton, C., Lucas, C., & Guest, E. (2024). *The operational risks of AI in large-scale biological attacks: Results of a red-team study* (RR-A2977-2). RAND Corporation. <https://doi.org/10.7249/RRA2977-2>
26. Narayanan, A., & Kapoor, S. (2024). *AI snake oil: What artificial intelligence can do, what it can't, and how to tell the difference*. Princeton University Press.
27. Ng, A. (2023). *Written statement before the U.S. Senate AI Insight Forum*.
28. NIST. (2023–2025). *AI 800-1 and AI risk management framework 1.0*.
29. NVIDIA. (2023). *NeMo guardrails*. EMNLP 2023.
30. Ojanen, A., Anttila, J., Thelitz, T. H. K., & Bjork, A. (2025). *Governing rapid technological change: Policy Delphi on the future of European AI governance*. arXiv. <https://doi.org/10.48550/arXiv.2512.15196>
31. Orben, A. (2020). The Sisyphean cycle of technology panics. *Perspectives on Psychological Science*, 15(5), 1143–1157. <https://doi.org/10.1177/1745691620919372>
32. OWASP. (2025, December). *Top 10 for agentic applications*.
33. Peppin, A., Reuel, A., Casper, S., Jones, E., Strait, A., Anwar, U., Agrawal, A., Kapoor, S., Koyejo, S., Pellat, M., Bommasani, R., Frosst, N., & Hooker, S. (2024). *The reality of AI and biorisk*. arXiv. <https://doi.org/10.48550/arXiv.2412.01946>
34. Shapira, N., Wendler, C., Yen, A., Sarti, G., Pal, K., Floody, O., Belfki, A., Loftus, A., Jannali, A. R., Prakash, N., Cui, J., Rogers, G., Brinkmann, J., Rager, C., Zur, A., Ripa, M.,

- Sankaranarayanan, A., Atkinson, D., Gandikota, R., Fiotto-Kaufman, J., Hwang, E., Orgad, H., Sahil, P.S., Taglicht, N., Shabtay, T., Ambus, A., Alon, N., Oron, S., Gordon-Tapiero, A., Kaplan, Y., Shwartz, V., Shaham, T. R., Riedl, C., Mirsky, R., Sap, M., Manheim, D., Ullman, T., & Bau, D. (2026). *Agents of chaos*. arXiv. <https://doi.org/10.48550/arXiv.2602.20021>
35. Sharma, M., Tong, M., Mu, J., Wei, J., Kruthoff, J., Goodfriend, S., Ong, E., Peng, A., Agarwal, R., Anil, C., Askill, A., Bailey, N., Benton, J., Bluemke, E., Bowman, S. R., Christiansen, E., Cunningham, H., Dau, A., Gopal, A., Gilson, R., Graham, L., Howard, L., Kalra, N., Lee, T., Lin, K., Lofgren, P., Mosconi, F., O'Hara, C., Olsson, C., Petrini, L., Rajani, S., Saxena, N., Silverstein, A., Singh, T., Summers, T., Tang, L., Troy, K. K., Weisser, C., Zhong, R., Zhou, G., Leike, J., Kaplan, J., & Perez, E. (2025). *Constitutional classifiers: Defending against universal jailbreaks across thousands of hours of red teaming*. arXiv. <https://doi.org/10.48550/arXiv.2501.18837>
 36. Tipton, E., Hallberg, K., Hedges, L. V., & Chan, W. (2016). Implications of Small Samples for Generalization: Adjustments and Rules of Thumb. *Evaluation Review*, 41(5), 472-505. <https://doi.org/10.1177/0193841X16655665>
 37. Weidinger, L., Mellor, J., Pegueroles, B. G., Marchal, N., Kumar, R., Lum, K., Akbulut, C., Diaz, M., Bergman, S., Rodriguez, M., Rieser, V., & Isaac, W. (2024). *STAR: SocioTechnical Approach to Red Teaming Language Models*. arXiv. <https://doi.org/10.48550/arXiv.2406.11757>
 38. Yin, R. K. (2013). Validity and generalization in future case study evaluations. *Evaluation*, 19(3), 321–332. <https://doi.org/10.1177/1356389013497081>