

Nicole Schneidman (SBN 319511)  
AI for Democracy Action Lab at  
PROTECT DEMOCRACY PROJECT  
P.O. Box 341423  
Los Angeles, CA 90034-9998  
(202) 579-4582  
nicole.schneidman@protectdemocracy.org

Deana K. El-Mallawany (SBN 674825)\*  
AI for Democracy Action Lab at  
PROTECT DEMOCRACY PROJECT  
15 Main Street, Suite 312  
Watertown, MA 02472  
(202) 579-4582  
deana.elmallaway@protectdemocracy.org

Ori Lev (SBN 452565)\*  
AI for Democracy Action Lab at  
PROTECT DEMOCRACY PROJECT  
2020 Pennsylvania Ave. NW, Suite #163  
Washington, D.C. 20006  
(202) 579-4582  
ori.lev@protectdemocracy.org

*Counsel for Amici Curiae*

*\*Application for admission pro hac vice  
forthcoming*

**IN THE UNITED STATES DISTRICT COURT  
FOR THE NORTHERN DISTRICT OF CALIFORNIA**

ANTHROPIC PBC,

*Plaintiff,*

v.

U.S. DEPARTMENT OF WAR, *et al.*,

*Defendants.*

Case No.: 3:26-cv-01996-RFL

**[PROPOSED] BRIEF OF *AMICI  
CURIAE* EMPLOYEES OF OPENAI  
AND GOOGLE IN THEIR PERSONAL  
CAPACITIES IN SUPPORT OF  
PLAINTIFF'S MOTION FOR A  
TEMPORARY RESTRAINING ORDER**

Judge: Hon. Rita F. Lin

Date:  
Time:  
Ctrm:

## TABLE OF CONTENTS

|                                                                                                                                |    |
|--------------------------------------------------------------------------------------------------------------------------------|----|
| TABLE OF CONTENTS.....                                                                                                         | 2  |
| TABLE OF AUTHORITIES .....                                                                                                     | 3  |
| IDENTITY AND INTEREST OF THE <i>AMICI</i> .....                                                                                | 1  |
| INTRODUCTION AND SUMMARY OF ARGUMENT .....                                                                                     | 3  |
| ARGUMENT .....                                                                                                                 | 4  |
| I.    The “Supply Chain Risk” Designation Is Improper Retaliation That Harms the Public Interest. ....                         | 4  |
| II.   The Concerns Underlying Anthropic’s “Red Lines” Are Real and Require a Response. ..                                      | 5  |
| III.  Mass Domestic Surveillance Powered by AI Poses Profound Risks to Democratic Governance — Even in Responsible Hands. .... | 6  |
| IV.   Fully Autonomous Lethal Weapons Systems Present Risks That Also Must Be Addressed. ....                                  | 9  |
| CONCLUSION.....                                                                                                                | 10 |

**TABLE OF AUTHORITIES**

**Cases**

*Hartman v. Moore*, 547 U.S. 250 (2006)..... 5

**Statutes**

10 U.S.C. § 3252..... 4

47 U.S.C. §1601..... 4

National Defense Authorization Act of 2026, Pub. L. 119-60, tit. XV, § 1532 ..... 5

Posse Comitatus Act, 18 U.S.C. § 1385 (originally enacted 1878) ..... 8

**Other Authorities**

Exec. Order No. 13873, 84 Fed. Reg. 22689 (May 17, 2019)..... 5

Jonathon W. Penney, *Chilling Effects: Online Surveillance and Wikipedia Use*, 31 Berkeley  
Tech. L.J. 117 (2016)..... 8

Privacy and Civil Liberties Oversight Board, *Report on the Telephone Records Program  
Conducted Under Section 215 of the USA PATRIOT Act* (Jan. 23, 2014)..... 8

Senate Select Committee to Study Governmental Operations with Respect to Intelligence  
Activities, Final Report, S. Rep. No. 94-755, Book III (1976)..... 8

U.S. Bureau of Labor Statistics, *Beyond the Numbers: Artificial Intelligence — Taking on a  
Bigger Role in Our Future Security* (May 3, 2021)..... 6

## IDENTITY AND INTEREST OF THE *AMICI*

*Amici* are engineers, researchers, scientists, and other professionals employed at U.S. frontier artificial intelligence laboratories. We build, train, and study the large-scale AI systems that serve a wide range of users and deployments, including in the consequential domains of national security, law enforcement, and military operations. We submit this brief not as spokespeople for any single company, but in our individual capacities as professionals with direct knowledge of what these systems can and cannot do, and what is at stake when their deployment outpaces the legal and ethical frameworks designed to govern them.

As a group, we are diverse in our politics and philosophies, but we are united in the conviction that today's frontier AI systems present risks when deployed to enable domestic mass surveillance or the operation of autonomous lethal weapons systems without human oversight, and that those risks require some kind of guardrails, whether via technical safeguards or usage restrictions. We view this conviction not as a result of any particular set of ideological or political commitments, but rather as a conclusion that follows from any reasonable evaluation of the capabilities and limitations of currently available frontier AI systems. It is this conviction that brings us before the Court to respectfully submit this brief, in the hopes that our understanding of the technology at issue, and our unique perspective as employees of companies currently engaged in fierce competition with Anthropic, will shed some light on the stakes of this case.

Individual *amici* are identified below. They sign the brief in their personal capacities only, and their companies and titles are included only to provide a sense of the perspectives they bring to this case.

- Grant Birkinbine, Member of Technical Staff – Security Engineer, OpenAI
- Anna-Luisa Brakman, Member of Technical Staff, OpenAI
- Sarah Cogan, Senior Software Engineer, Google DeepMind
- Jeff Dean, Chief Scientist, Google
- Michael Dennis, Senior Research Scientist, Google DeepMind
- Sanjeev Dhanda, Senior Staff Software Engineer, Google DeepMind

- Rasmi Elasmr, Senior Research Engineer, Google DeepMind
- Brian Fioca, Member of Go to Market Staff, OpenAI
- Aaron Friel, Member of Technical Staff – Software Engineer, OpenAI
- Leo Gao, Member of Technical Staff, OpenAI
- Edward Grefenstette, Director of Research, Google DeepMind
- Alexander Irpan, Research Scientist, Google DeepMind
- Rishub Jain, Research Engineer, Google DeepMind
- Manas Joglekar, Member of Technical Staff, OpenAI
- Kathy Korevec, Director, Product at Google
- Shrinu Kushagra, Research Scientist, Google
- Teddy Lee, Member of GTM Staff, OpenAI
- Sharon Lin, Research Engineer, Google DeepMind
- Soukaina Mansour, Creative Community Lead, OpenAI
- Ian McKenzie, Research Engineer, Google DeepMind
- Pamela Mishkin, Research, OpenAI
- Roman Novak, Research Scientist, OpenAI
- Zach Parent, Forward Deployed Engineer, OpenAI
- Andrew Schmidt, Model Designer, OpenAI
- Noah Siegel, Senior Research Engineer, Google DeepMind
- Jordan Sitkin, Member of Technical Staff, OpenAI
- Chang Sun, Member of Data Science Staff, OpenAI
- Sean Talts, Staff Software Engineer, Google
- Alexander Matt Turner, Research Scientist, Google DeepMind
- Anna Wang, Research Scientist, Google DeepMind
- Zhengdong Wang, Senior Research Engineer, Google DeepMind
- Jonathan Ward, Member of Technical Staff, OpenAI
- Jason Wolfe, Member of Technical Staff, OpenAI

- Kate Woolverton, Senior Software Engineer, Google DeepMind
- Gabriel Wu, Member of Technical Staff – Research Engineer, OpenAI
- Cathy Yeh, Member of Technical Staff, OpenAI
- Jelle Zijlstra, Member of Technical Staff, OpenAI

## INTRODUCTION AND SUMMARY OF ARGUMENT

This case arises from the Pentagon delivering on its threat to designate Anthropic a “supply chain risk” if the company declined to agree to remove limitations on the use of its AI systems for domestic mass surveillance or fully autonomous lethal weapons systems. If it were no longer satisfied with the agreed-upon terms of its contract with Anthropic, the Defendants could have simply canceled the contract and purchased the services of another leading AI company. Instead, Defendants recklessly invoked national security authorities intended to protect the procurement process from interference by foreign adversaries. If allowed to proceed, this effort to punish one of the leading U.S. AI companies will undoubtedly have consequences for the United States’ industrial and scientific competitiveness in the field of artificial intelligence and beyond. And it will chill open deliberation in our field about the risks and benefits of today’s AI systems. Because we understand the risks of frontier AI systems and the need for guardrails, and because we believe that speaking openly about them is of paramount importance, we submit this brief.

We offer three arguments.

First, the government’s designation of Anthropic as a supply chain risk was an improper and arbitrary use of power that has serious ramifications for our industry. While we are not privy to the details of how Anthropic and the Pentagon’s contractual relationship broke down, we are concerned that the Defendants’ action harms public debate on the risks and benefits of AI as well as U.S. competitiveness in the field of AI and innovation more broadly.

Second, the technical concerns animating Anthropic’s “red lines” are legitimate and widely recognized within our scientific community as requiring some kind of response. The best

1 currently available AI systems cannot safely or reliably handle fully autonomous lethal targeting,  
 2 and should not be available for domestic mass surveillance of the American people. While there  
 3 are various ways to establish these guardrails, we agree that these guardrails must be in place.

4 Third, as AI professionals, we understand that the substantive risks of the two use cases  
 5 at issue are profound. AI-enabled mass domestic surveillance would transform the fragmented  
 6 data ecosystem that already surrounds American life into a unified, real-time instrument for  
 7 monitoring the entire population. Even the awareness that such capability exists creates a chilling  
 8 effect on democratic participation. Autonomous lethal weapons systems, as currently designed  
 9 and deployed, cannot reliably distinguish combatants from civilians, cannot explain their  
 10 targeting decisions, and cannot engage in human accountability structures. These concerns  
 11 require a response.

12 For these reasons, we urge the Court to grant the relief requested by Anthropic.

### 13 **ARGUMENT**

#### 14 **I. The “Supply Chain Risk” Designation Is Improper Retaliation That Harms the** 15 **Public Interest.**

16 This case poses a question of seismic importance for our industry, our national security,  
 17 and our democracy: What happens when the government uses its national security authorities to  
 18 punish a private company for maintaining safeguards on certain uses of its AI systems while  
 19 speaking to why those safeguards exist and why they matter?

20 In early March 2026, the Pentagon officially designated Anthropic as a supply chain risk,  
 21 following earlier threats to do so. While we are not privy to the details of their negotiations, the  
 22 Defendants had the option simply to drop Anthropic’s contract if it no longer wished to be bound  
 23 by its terms. The supply chain risk designation is a mechanism for excluding from the defense  
 24 industrial base vendors who pose a genuine threat to the integrity of military systems.<sup>1</sup> It is  
 25 scarcely used, and then for foreign adversary-controlled companies, compromised suppliers, and  
 26

---

27 <sup>1</sup> See 10 U.S.C. § 3252 (supply chain security authority); *see also* Secure and Trusted Networks  
 28 Security Act of 2019, 47 U.S.C. §1601 *et seq.* (requiring the Federal Communications  
 Commission to publish list of prohibited equipment).

1 contractors whose products create exploitable vulnerabilities.<sup>2</sup> Anthropic is a domestic AI  
 2 developer<sup>3</sup> that has worked with the Pentagon on military applications of AI systems since last  
 3 year.

4 The Pentagon's decision to reach for supply chain risk authority in response to  
 5 Anthropic's contract negotiations introduces an unpredictability in our industry that undermines  
 6 American innovation and competitiveness. It chills professional debate on the benefits and risks  
 7 of frontier AI systems and various ways that risks can be addressed to optimize the technology's  
 8 deployment. The United States' thriving AI ecosystem leads the rest of the world largely due to  
 9 the competition and flow of ideas between different AI companies. By silencing one lab, the  
 10 government reduces the industry's potential to innovate solutions. The resulting harm has  
 11 constitutional dimensions as well, undermining the freedom to engage in public debate about  
 12 how powerful technologies should be governed. *See Hartman v. Moore*, 547 U.S. 250, 256  
 13 (2006) ("[T]he First Amendment prohibits government officials from subjecting an individual to  
 14 retaliatory actions . . . for speaking out.").

## 15 **II. The Concerns Underlying Anthropic's "Red Lines" Are Real and Require a** 16 **Response.**

17 As AI professionals, we recognize that frontier AI is a powerful technology that could  
 18 have many benefits for humanity but also carries many risks. The risks are not hypothetical.  
 19 They are structural. They follow from the nature of the technology itself, at least as it exists  
 20 today, and from what happens when institutions, however well-intentioned, acquire capabilities  
 21 that exceed the oversight mechanisms designed to check them.

22 That is why it is important to put guardrails around the domains in which these systems  
 23 carry intolerable risk as they are currently constructed. A child's tricycle can physically be  
 24 driven on an interstate, but we do not allow it because of the risks of using the technology in that

---

25 <sup>2</sup> *See, e.g.*, Exec. Order No. 13873, 84 Fed. Reg. 22689 (May 17, 2019) (authorizing import bans  
 26 on technology threats under emergency trade authorities, aimed at Chinese vendors Huawei and ZTE).

27 <sup>3</sup> It is at least notable that Congress has already largely prohibited acquisition of AI systems from  
 28 China, Russia, North Korea, and Iran. National Defense Authorization Act of 2026, Pub. L. 119-  
 60, tit. XV, § 1532. There is no equivalent provision for U.S.-based systems providers. *Id.*

environment. Mass domestic surveillance and autonomous lethal weapons systems are the equivalently reckless domain for today's frontier models. The considered judgment, shared widely across the AI development community, is that these applications of current AI technology carry risks so severe, and threaten harm so impossible to repair after the fact, that some kind of guardrails — whether contractual or technical — are necessary to constrain them in the absence of robust, genuinely effective governance frameworks. For a system vendor to insist that those boundaries be honored as a condition of access to its software is not arbitrary, anticompetitive, or contrary to the public interest.

The legal vacuum in which these contractual terms exist makes them only more important. The United States currently has no comprehensive federal law governing the use of AI by military or intelligence agencies in domestic contexts. There is no statutory framework requiring transparency, judicial oversight, or meaningful accountability for AI-driven surveillance at scale. There is no enforceable legal standard governing when an autonomous weapons system may select and engage a target. In the absence of public law, the contractual and technological requirements that AI developers impose on the use of their systems represent a vital safeguard against their catastrophic misuse.

### **III. Mass Domestic Surveillance Powered by AI Poses Profound Risks to Democratic Governance — Even in Responsible Hands.**

The risks of AI-enabled mass domestic surveillance merit greater public understanding. At its core, AI-enabled mass surveillance means the ability to monitor, analyze, and act on the behavior of an entire population continuously and in real time. The devices and data streams required to do this already exist. As of 2018,<sup>4</sup> there were approximately 70 million surveillance cameras operating in the United States across airports, subway stations, parking lots, storefronts, and street corners. Every smartphone continuously broadcasts location data to carriers and dozens of applications. Credit and debit cards generate a timestamped record of nearly every

---

<sup>4</sup> U.S. Bureau of Labor Statistics, *Beyond the Numbers: Artificial Intelligence — Taking on a Bigger Role in Our Future Security* (May 3, 2021), <https://www.bls.gov/opub/btn/volume-10/investigation-and-security-services.htm>.

1 commercial transaction Americans make. Social media platforms log not just what people post,  
2 but what they read, how long they browse, and what they posted before deleting it. Employers,  
3 insurers, and data brokers have assembled behavioral profiles on most American adults that are  
4 already, in many cases, available for government purchase without a warrant. What does not yet  
5 exist is the AI layer that transforms this sprawling, fragmented data landscape into a unified,  
6 real-time surveillance apparatus. Today, these streams are siloed, inconsistent, and require  
7 significant human effort to connect. From our vantage point at frontier AI labs, we understand  
8 that an AI system used for mass surveillance could dissolve those silos, correlating face  
9 recognition data with location history, transaction records, social graphs, and behavioral patterns  
10 across hundreds of millions of people simultaneously.

11       The mere existence of such a capability in government hands — even if never activated  
12 against a specific individual — changes the character of public life in a democracy. Behavioral  
13 scientists and legal scholars have long documented what is sometimes called the “panopticon  
14 effect”: when people believe they may be observed, they modify their behavior as if they are  
15 always being observed, regardless of whether anyone is actually watching. The journalist thinks  
16 twice before calling a source inside the military, knowing the call could be logged and cross-  
17 referenced. The activist softens her public messaging, calculating that visibility now carries risk  
18 it didn’t carry before. The academic researcher avoids certain search terms — not because the  
19 research is wrong, but because she doesn’t want to surface in a database. None of these people  
20 have been targeted. None have been punished. But their behavior has already been constrained,  
21 and with it the democratic functions they serve — a free press, political organizing, open  
22 intellectual inquiry — have been quietly degraded. These chilling effects require no abuse, only  
23 the awareness that the capability exists.

24       History offers ample warning. The FBI’s COINTELPRO program, which ran from 1956  
25 to 1971 and was exposed years later, demonstrated how domestic intelligence powers justified by  
26 security concerns were systematically turned against civil rights leaders, journalists, and political  
27 dissidents. The program did not merely surveil its targets. It fabricated evidence, sent anonymous  
28

1 letters designed to destroy marriages and careers, tipped off employers, and worked to discredit  
 2 Martin Luther King, Jr. after he was awarded the Nobel Peace Prize.<sup>5</sup> It operated for fifteen years  
 3 before Congress learned of its existence. AI does not merely replicate those dangers — it  
 4 multiplies them by orders of magnitude, automating at national scale what previously required  
 5 hundreds of human operatives.

6 Further enhancing the risk terrain for AI's deployment in this context, the Pentagon  
 7 operates under a legal framework oriented toward external threats and warfighting, not domestic  
 8 civil life. The Posse Comitatus Act, passed in 1878 in direct response to the use of federal troops  
 9 to police American civilians during Reconstruction, reflects a constitutional tradition of keeping  
 10 military power categorically separate from domestic governance.<sup>6</sup> When the Pentagon acts  
 11 domestically, it is operating in legal territory it was not designed for, with oversight structures  
 12 that were not built to catch domestic abuses. That is in part why the bulk data collection  
 13 programs by the Pentagon's own National Security Agency (NSA), revealed by Edward  
 14 Snowden in 2013, were so shocking and produced measurable chilling effects on lawful speech  
 15 and inquiry. A study published in the Berkeley Technology Law Journal found statistically  
 16 significant drops in traffic to Wikipedia articles on terrorism-related topics following the  
 17 Snowden revelations, likely as ordinary people adjusted their online behavior in response to  
 18 awareness that their searches were potentially being monitored.<sup>7</sup>

19 The harms from building this infrastructure are not easily undone, as we understand in  
 20 our field. Data collected on a population does not expire. A database of location records,  
 21 behavioral profiles, and social graphs built today will still exist years from now, accessible to

---

22 <sup>5</sup> See Senate Select Committee to Study Governmental Operations with Respect to Intelligence  
 23 Activities, Final Report, S. Rep. No. 94-755, Book III, at 3, 8, 26, 138, 143-144 (1976).

24 <sup>6</sup> Posse Comitatus Act, 18 U.S.C. § 1385 (originally enacted 1878) (prohibiting use of the  
 25 military to execute civilian laws except as expressly authorized by the Constitution or Act of  
 Congress).

26 <sup>7</sup> Jonathon W. Penney, *Chilling Effects: Online Surveillance and Wikipedia Use*, 31 Berkeley  
 27 Tech. L.J. 117 (2016). For more on NSA program's chilling effects, see generally Privacy and  
 28 Civil Liberties Oversight Board, *Report on the Telephone Records Program Conducted Under  
 Section 215 of the USA PATRIOT Act* (Jan. 23, 2014),  
[https://documents.pclob.gov/prod/Documents/OversightReport/ec542143-1079-424a-84b3-  
 acc354698560/215-Report\\_on\\_the\\_Telephone\\_Records\\_Program.pdf](https://documents.pclob.gov/prod/Documents/OversightReport/ec542143-1079-424a-84b3-acc354698560/215-Report_on_the_Telephone_Records_Program.pdf).

1 whoever controls it under whatever political conditions prevail then. That data would feed into  
2 an AI-powered surveillance infrastructure that, once constructed, tends to expand rather than  
3 contract. Agencies find new uses for existing capabilities, authorities get quietly reinterpreted,  
4 and the political cost of dismantling something already built is almost always higher than the  
5 cost of letting it continue and grow. One lesson of the Snowden revelations is that technology  
6 built for international espionage has a way of being used for domestic surveillance without clear  
7 legal boundaries. The boundary the Posse Comitatus Act was designed to protect, once eroded by  
8 the establishment of a Pentagon-controlled domestic surveillance apparatus, may prove  
9 practically impossible to restore.

10 We do not suggest that the Defendants intend to misuse such capabilities. We suggest  
11 that the question of intent is the wrong question. Democratic governance does not rest on the  
12 good intentions of those in power. It rests on structural constraints that make abuse difficult  
13 regardless of intent. AI-enabled mass domestic surveillance, deployed without transparent legal  
14 constraints and independent oversight, removes those structural protections in ways that no  
15 amount of good faith can replace.

#### 16 **IV. Fully Autonomous Lethal Weapons Systems Present Risks That Also Must Be** 17 **Addressed.**

18 As professionals at frontier AI companies, we also recognize widely shared concerns  
19 around the deployment of lethal autonomous weapons systems. Current AI models are not  
20 reliable enough to bear the responsibility of making lethal targeting decisions entirely alone, and  
21 the risks of their deployment for that purpose require some kind of response and guardrails.

22 Lethal autonomous weapons systems are no longer hypothetical. They are already being  
23 deployed with decreasing levels of human involvement, and their failures are already  
24 documented. Our experience convinces us that current AI systems have limitations, as pattern-  
25 matching systems trained on historical data, that create an unacceptable risk of deployment in  
26 fully autonomous forms. This is because, while expert pattern-matching systems perform well in  
27 conditions that resemble their training environment, they have a significant potential to degrade  
28

1 in novel, ambiguous conditions. They cannot be trusted to identify targets with perfect accuracy,  
2 and they are incapable of making the subtle contextual tradeoffs between achieving an objective  
3 and accounting for collateral effects that a human can. While AI systems can assemble and  
4 evaluate information quickly and provide valuable information to human decisionmakers, they  
5 also have the potential to hallucinate, meaning that a human must be able to confirm the  
6 accuracy of the critical information before a lethal munition is launched at a human target. Their  
7 chain of reasoning is often hidden from their operators, and their internal workings are opaque  
8 even to their developers. And the decisions they make in lethal contexts are irreversible. Even if  
9 fully autonomous weapons systems are inevitable, they cannot be safely deployed without some  
10 kind of guardrails to make their use reasonable. That is a technical judgment, not a political one.

### 11 CONCLUSION

12 The government has legitimate interests in ensuring that AI capabilities are available to  
13 serve national security. But national security is not served by reckless designations of the  
14 military’s American technology partners as a “supply chain risk” or the suppression of public  
15 discourse on AI safety. Nor is the United States’ competitiveness in AI development served by  
16 the Defendants’ retaliation against one of the leading American companies in our field. Until a  
17 legal framework exists to contain the risks of deploying frontier AI systems, the ethical  
18 commitments of AI developers — and their willingness to defend those commitments publicly  
19 — are not obstacles to good governance or innovation. They are contributions to it. The Court  
20 should say so.

1 Dated: March 9, 2026

/s/ Nicole Schneidman

Nicole Schneidman (SBN 319511)  
AI for Democracy Action Lab at  
PROTECT DEMOCRACY PROJECT  
P.O. Box 341423  
Los Angeles, CA 90034-9998  
(202) 579-4582  
nicole.schneidman@protectdemocracy.org

6 Ori Lev (SBN 452565)\*  
7 AI for Democracy Action Lab at  
8 PROTECT DEMOCRACY PROJECT  
9 2020 Pennsylvania Ave. NW, Suite #163  
10 Washington, D.C. 20006  
11 (202) 579-4582  
12 ori.lev@protectdemocracy.org

Deana K. El-Mallawany (SBN 674825)\*  
AI for Democracy Action Lab at  
PROTECT DEMOCRACY PROJECT  
15 Main Street, Suite 312  
Watertown, MA 02472  
(202) 579-4582  
deana.elmallaway@protectdemocracy.org

*Counsel for Amici Curiae*

*\*Application for admission pro hac vice  
forthcoming*