

MICHAEL J. MONGAN (SBN 250374)
 michael.mongan@wilmerhale.com
 WILMER CUTLER PICKERING
 HALE AND DORR LLP
 50 California Street, Suite 3600
 San Francisco, CA 94111
 Telephone: (628) 235-1000

EMILY BARNET (*pro hac vice*)
 emily.barnet@wilmerhale.com
 WILMER CUTLER PICKERING
 HALE AND DORR LLP
 7 World Trade Center
 250 Greenwich Street
 New York, NY 10007
 Telephone: (212) 230-8800

Attorneys for Plaintiff Anthropic PBC

KELLY P. DUNBAR (*pro hac vice*)
 kelly.dunbar@wilmerhale.com
 JOSHUA A. GELTZER (*pro hac vice*)
 joshua.geltzer@wilmerhale.com
 KEVIN M. LAMB (*pro hac vice*)
 kevin.lamb@wilmerhale.com
 SUSAN HENNESSEY (*pro hac vice*)
 susan.hennessey@wilmerhale.com
 LAUREN MOXLEY BEATTY (SBN 308333)
 lauren.beatty@wilmerhale.com
 MATTHEW E. MORRIS (*pro hac vice*)
 matt.morris@wilmerhale.com
 BARDIA VASEGHI (*pro hac vice*)
 bardia.vaseghi@wilmerhale.com
 DAKOTA C. FOSTER (*pro hac vice*)
 dakota.foster@wilmerhale.com
 WILMER CUTLER PICKERING
 HALE AND DORR LLP
 2100 Pennsylvania Avenue NW
 Washington, DC 20037
 Telephone: (202) 663-6000

**UNITED STATES DISTRICT COURT
 NORTHERN DISTRICT OF CALIFORNIA**

ANTHROPIC PBC,

Plaintiff,

v.

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

Defendants.

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

**ANTHROPIC PBC'S MOTION FOR
 SUMMARY JUDGMENT**

Judge: Hon. Rita F. Lin
 Hearing Date: July 30, 2026
 Time: 10:00 a.m. PT

TABLE OF CONTENTS

1		
2	INTRODUCTION	1
3	BACKGROUND	2
4	A. Anthropic And Its Mission To Responsibly Develop AI Tools	2
5	B. Anthropic’s Partnership With The Federal Government	2
6	C. The Present Dispute	3
7	ARGUMENT	6
8	I. Anthropic Is Entitled To Summary Judgment On Its First Amendment	
9	Claim Because Defendants Retaliated Against Anthropic For Its	
	Protected Expression	6
10	A. The Undisputed Facts Establish A Prima Facie Case Of Retaliation	7
11	B. Defendants Cannot Rebut The Prima Facie Case	10
12	II. Secretary Hegseth’s Actions Violate The APA Because They Are	
13	Arbitrary, Capricious, And Contrary To § 3252’s Procedural And	
	Substantive Requirements	12
14	A. Secretary Hegseth’s February 27 Directive Is Final And Unlawful	
15	Agency Action	12
16	B. The March Determination Is Independently Unlawful	13
17	C. Additional Materials Properly Considered Here Reinforce The APA	
	Violations	20
18	III. Defendants Violated The Due Process Clause By Blacklisting Anthropic	
19	Without Notice Or An Opportunity To Be Heard	21
20	IV. The Presidential Directive Violates The Separation Of Powers	24
21	V. Implementation Of The Presidential Directive Violated 5 U.S.C. § 558	24
22	CONCLUSION	25

TABLE OF AUTHORITIES

Page(s)

CASES

<i>Al Haramain Islamic Found. v. Dep't of Treasury</i> , 686 F.3d 965 (9th Cir. 2012)	21, 23, 24
<i>Alpha Energy Savers, Inc. v. Hansen</i> , 381 F.3d 917 (9th Cir. 2004)	7, 9
<i>Am. Bus Ass'n v. Slater</i> , 231 F.3d 1 (D.C. Cir. 2000)	24
<i>Am. Mfrs. Mut. Ins. Co. v. Sullivan</i> , 526 U.S. 40 (1999)	21
<i>Am. Textile Mfrs. Inst., Inc. v. Donovan</i> , 452 U.S. 490 (1981)	19
<i>Am. Fed'n of Gov't Employees, AFL-CIO v. Trump</i> , 167 F.4th 1247 (9th Cir. 2026)	11, 12
<i>Anderson v. Liberty Lobby, Inc.</i> , 477 U.S. 242 (1986)	6
<i>Ariz. Students' Ass'n v. Ariz. Bd. of Regents</i> , 824 F.3d 858 (9th Cir. 2016)	6
<i>Bennett v. Spear</i> , 520 U.S. 154 (1997)	12
<i>Burlington Truck Lines Inc. v. United States</i> , 371 U.S. 156 (1962)	18
<i>Burwell v. Hobby Lobby Stores, Inc.</i> , 573 U.S. 682 (2014)	7
<i>Cal. Motor Transp. Co. v. Trucking Unlimited</i> , 404 U.S. 508 (1972)	8
<i>CarePartners, LLC v. Lashway</i> , 545 F.3d 867 (9th Cir. 2008)	10
<i>Cleveland Bd. of Educ. v. Loudermill</i> , 470 U.S. 532 (1985)	23
<i>Compassion Over Killing v. FDA</i> , 849 F.3d 849 (9th Cir. 2017)	6
<i>Conner v. City of Santa Ana</i> , 897 F.2d 1487 (9th Cir. 1990)	21, 22

1	<i>Coszalter v. City of Salem,</i>	
2	320 F.3d 968 (9th Cir. 2003)	9
3	<i>Dep't of Com. v. New York,</i>	
4	588 U.S. 752 (2019).....	19, 20
5	<i>DHS v. Regents of the Univ. of Cal.,</i>	
6	591 U.S. 1 (2020).....	13, 19
7	<i>Fikre v. FBI,</i>	
8	35 F.4th 762 (9th Cir. 2022)	21
9	<i>Fuentes v. Shevin,</i>	
10	407 U.S. 67 (1972).....	23
11	<i>Garza v. Woods,</i>	
12	150 F.4th 1118 (9th Cir. 2025)	23
13	<i>Greene v. McElroy,</i>	
14	360 U.S. 474 (1959).....	21
15	<i>Hartman v. Moore,</i>	
16	547 U.S. 250 (2006).....	6
17	<i>Janus v. Am. Fed'n of State, Cty., & Mun. Emps., Council 31,</i>	
18	585 U.S. 878 (2018).....	7, 8
19	<i>Kartseva v. Dep't of State,</i>	
20	37 F.3d 1524 (D.C. Cir. 1994).....	21
21	<i>Kennedy v. Braidwood Mgmt., Inc.,</i>	
22	606 U.S. 748 (2025).....	16, 17
23	<i>Lands Council v. Powell,</i>	
24	395 F.3d 1019 (9th Cir. 2005)	20
25	<i>Mathews v. Eldridge,</i>	
26	424 U.S. 319 (1976).....	22, 24
27	<i>Mills v. Alabama,</i>	
28	384 U.S. 214 (1966).....	9
	<i>Nat. Res. Def. Council, Inc. v. Pritzker,</i>	
	828 F.3d 1125 (9th Cir. 2016)	14
	<i>Nat'l Ass'n of Home Builders v. Norton,</i>	
	340 F.3d 835 (9th Cir. 2003)	14
	<i>Nat'l Council of Resistance of Iran v. Dep't of State,</i>	
	51 F.3d 192 (D.C. Cir. 2001).....	23
	<i>New York Times Co. v. Sullivan,</i>	
	376 U.S. 254 (1964).....	8

1	<i>People’s Mojahedin Org. of Iran v. Dep’t of State</i> ,	
2	613 F.3d 220 (D.C. Cir. 2010).....	23
3	<i>Ralls Corp. v. CFIUS</i> ,	
4	758 F.3d 296 (D.C. Cir. 2014).....	21
5	<i>Riley’s Am. Heritage Farms v. Elsasser</i> ,	
6	32 F.4th 707 (9th Cir. 2022)	11
7	<i>Scott v. Harris</i> ,	
8	550 U.S. 372 (2007).....	10
9	<i>Snyder v. Phelps</i> ,	
10	562 U.S. 443 (2011).....	7, 8
11	<i>State v. Su</i> ,	
12	121 F.4th 1 (9th Cir. 2024)	24
13	<i>Thomas v. County of Riverside</i> ,	
14	763 F.3d 1167 (9th Cir. 2014)	10
15	<i>Thompson v. U.S. Dep’t of Lab.</i> ,	
16	885 F.2d 551 (9th Cir. 1989)	20
17	<i>Trifax Corp. v. Dist. of Columbia</i> ,	
18	314 F.3d 641 (D.C. Cir. 2003).....	21
19	<i>U.S. Army Corps of Eng’rs v. Hawkes Co.</i> ,	
20	578 U.S. 590 (2016).....	12
21	<i>Ulrich v. City & Cty. of San Francisco</i> ,	
22	308 F.3d 968 (9th Cir. 2002)	21
23	<i>United States ex rel. Accardi v. Shaughnessy</i> ,	
24	347 U.S. 260 (1954).....	14
25	<i>White v. Lee</i> ,	
26	227 F.3d 1214 (9th Cir. 2000)	10
27	<i>Youngstown Sheet & Tube Co. v. Sawyer</i> ,	
28	343 U.S. 579 (1952).....	24

STATUTES, RULES, AND REGULATIONS

5 U.S.C. § 558.....	24, 25
5 U.S.C. § 706.....	<i>passim</i>
10 U.S.C. § 3203	24
10 U.S.C. § 3204.....	24
10 U.S.C. § 3252.....	<i>passim</i>
41 U.S.C. § 1327.....	5

1	41 U.S.C. § 3303	24
2	41 U.S.C. § 3304	24
3	41 U.S.C. § 4713	5
4	48 C.F.R. § 9.402	24
5	48 C.F.R. § 9.406-2	24
6	48 C.F.R. § 9.407-2	24
7	48 C.F.R. § 49.101	11
8	48 C.F.R. § 49.402-3	11
9	48 C.F.R. § 52.212-4	11
10	48 C.F.R. § 52.243-4	11
11	48 C.F.R. § 52.249-1	11
12	48 C.F.R. § 52.249-2	11
13	48 C.F.R. § 52.249-3	11
14	48 C.F.R. § 52.249-4	11
15	48 C.F.R. § 52.249-5	11
16	48 C.F.R. § 52.249-6	11
17	48 C.F.R. § 239.7304	14

OTHER AUTHORITIES

19	S. Rep. No. 111-201 (2010)	16
----	----------------------------------	----

NOTICE OF MOTION AND MOTION

PLEASE TAKE NOTICE that on July 30, 2026 at 10:00 a.m. PT, Plaintiff Anthropic PBC will and hereby does move the Court for an order granting Anthropic's Motion for Summary Judgment. Anthropic's motion is based on this Notice of Motion, the supporting Memorandum, and the Declarations filed herewith. Anthropic requests that the Court grant Anthropic summary judgment on each of its claims.

MEMORANDUM OF POINTS AND AUTHORITIES

INTRODUCTION

The undisputed facts establish that Defendants' campaign of retaliation against Anthropic violated the Constitution and the Administrative Procedure Act.

Defendants have been remarkably transparent about what happened here. They view Anthropic as providing the best frontier AI model and engaged the company to help the Department of War and other agencies protect America’s national security. They acknowledge that Anthropic has always made its models available subject to certain use restrictions, consistent with its longstanding advocacy for safe and responsible use of AI. In February 2026, Defendants insisted that Anthropic abandon two contractual restrictions that prohibited them from using the model for lethal autonomous warfare or mass surveillance of Americans. When Anthropic respectfully declined to abandon those restrictions, and publicly emphasized their importance, Defendants abruptly decided to punish the company. They barred every federal agency and every military contractor from doing business with Anthropic and designated the company a threat to national security.

Everyone agrees that the Executive Branch can lawfully decide to stop contracting with a vendor that declines to accept its preferred contract terms. But the Executive Branch remains subject to the requirements of the Constitution and federal law. It may not retaliate against an American company in violation of the First Amendment; it may not disregard procedural and substantive requirements imposed by Congress; it may not exercise its statutory national-security authorities in a manner that is arbitrary and capricious; and it may not punish a company without providing notice and an opportunity to respond.

As the undisputed record shows, Defendants did each of those things here. Their campaign of retaliation commenced with statements from the President and the Secretary of War that acknowledged—on their face—that the Executive was taking action against Anthropic for its “rhetoric” and its “ideology.” The Secretary designated Anthropic a “supply chain risk” under 10 U.S.C. § 3252 without taking the procedural steps required by Congress and based on rationales that exceed the statutory definition of a supply chain risk—and that the record shows to be arbitrary, false, and pretextual. Each of the unprecedented sanctions imposed by Defendants occurred without any pre-deprivation process, in violation of the Due Process Clause. On this record, Anthropic is entitled to judgment as a matter of law.

BACKGROUND

A. Anthropic And Its Mission To Responsibly Develop AI Tools

Anthropic is a frontier AI company founded in 2021. It offers an industry-leading large language model (LLM), known as Claude, which is capable of interpreting and responding to a wide range of user inputs in an intelligent, human-like manner. Kaplan Decl. ¶ 16.¹ LLMs are algorithmic systems that acquire predictive capabilities, allowing them to perform complex tasks in a fraction of the time it would take humans. *Id.* ¶ 15.

Anthropic was founded on the animating principle that the most capable AI systems should also be the safest. *Id.* ¶¶ 8-11. That principle is embedded in Anthropic’s corporate charter and governance. *Id.* ¶ 8. As a public benefit corporation, Anthropic balances stockholder interests with its public benefit purpose of responsibly developing and maintaining advanced AI tools for the long-term benefit of humanity. *See id.* ¶¶ 8-9.

B. Anthropic’s Partnership With The Federal Government

Anthropic began commercial deployment of Claude in 2023, *id.* ¶ 12, and started working with the federal government the following year, Ramasamy Decl. ¶ 7. In November 2024, Anthropic began working with a defense technology partner to provide AI-enabled capabilities to U.S. intelligence and defense agencies. *Id.* ¶ 8. On its own initiative and at its own expense,

¹ All declarations cited in this motion are attached and have been filed concurrently.

1 Anthropic built a new version of Claude (known as “Claude Gov”) that was tailored to national-
2 security needs and classified environments. *Id.* ¶¶ 31-32.

3 Every Anthropic model is subject to a Usage Policy that prohibits certain uses of Claude.
4 *Id.* ¶ 25. Claude Gov is subject to a government-specific addendum to the Usage Policy that
5 removes certain restrictions for national security customers that would otherwise apply to civilian
6 users. Kaplan Decl. ¶¶ 27-28. From the outset of the Department’s partnership with Anthropic, the
7 Department has operated subject to the Usage Policy and a government-specific addendum. *Id.*
8 ¶¶ 28-29; Ramasamy Decl. ¶ 26.

9 Since March 2025, the Department has accessed Claude Gov through classified cloud
10 environments; entrusted Anthropic to handle classified national security data; and awarded the
11 company a Top Secret facility clearance. Ramasamy Decl. ¶¶ 33, 58. In the summer of 2025, the
12 Department awarded Anthropic a two-year contract worth up to \$200 million. *Id.* ¶¶ 9, 69. Until
13 the events that gave rise to this litigation, Claude had reportedly become the Department’s most
14 widely deployed frontier AI model. Kaplan Decl. ¶ 26. According to the Department, Claude was
15 “far and away [its] best model.” Heck Decl. ¶ 9.

16 Other agencies have also made extensive use of Claude and praised it. In 2025, the General
17 Services Administration (GSA) announced a first-of-its-kind agreement with Anthropic to deliver
18 Claude across all three branches of government for just \$1 per agency. Ramasamy Decl. ¶ 71. And
19 the General Counsel of the Department of Veterans Affairs, for example, called Claude and
20 related tools “incredible” and pledged to “make full use of” them. AR 537.

21 **C. The Present Dispute**

22 In fall 2025, the Department and Anthropic began discussing a new deployment of Claude
23 on the Department’s “GenAI.mil” platform. Kaplan Decl. ¶ 32. For the first time, the Department
24 demanded that Anthropic discard its Usage Policy and agree to a new contractual provision
25 allowing the Department to use Claude “for all lawful uses.” *Id.* Anthropic agreed to
26 accommodate the Department by significantly revising its usage restrictions. *Id.* ¶ 33; AR 7. But it
27 respectfully declined to discard two narrow restrictions grounded in its founding commitments: it
28 would not authorize use of Claude for lethal autonomous warfare or mass surveillance of

1 Americans. Kaplan Decl. ¶ 33; AR 1-2. Anthropic’s CEO told the Under Secretary of War for
2 Research and Engineering that the company would assist in offboarding Claude from the
3 Department’s systems if it concluded Anthropic was not the right vendor for its needs. AR 2, 29.

4 On February 24, Anthropic’s leadership met with Secretary Hegseth and his team at the
5 Pentagon. Heck Decl. ¶ 13. The Secretary did not say anything about Anthropic’s AI models
6 being unsafe, insecure, or subject to compromise. *Id.* ¶ 18. To the contrary, he praised Claude’s
7 “exquisite capabilities” and told Anthropic that the Department “would love to work with” the
8 company. *Id.* ¶¶ 13, 16. But he insisted on Anthropic removing its two remaining restrictions. *Id.*
9 ¶¶ 13, 17. And he issued an ultimatum: Unless Anthropic agreed to the “all lawful uses” contract
10 term by 5 p.m. ET on February 27, he would either designate the company as a “supply chain
11 risk” *or* invoke the Defense Production Act to compel Anthropic to comply with his demand. *Id.*
12 ¶ 17. On the eve of that deadline, Anthropic’s CEO publicly explained why Anthropic could not
13 “in good conscience” abandon its commitments to AI safety. *See* Mongan Decl. Ex. 1 at 1-2.

14 The next day, before the Secretary’s 5 p.m. deadline, President Trump issued a statement
15 on social media (the “Presidential Directive”). He branded Anthropic a “Radical Left AI
16 company” that threatened “AMERICAN LIVES” and “direct[ed] EVERY Federal Agency in the
17 United States Government to IMMEDIATELY CEASE all use of Anthropic’s technology.” AR
18 255A. Federal agencies moved quickly to carry out the Presidential Directive. That same day,
19 GSA removed Anthropic from the government’s primary procurement marketplace, AR 394-95,
20 and the Department of Health and Human Services disabled access to Claude, AR 335. Within
21 days, the Departments of State, Treasury, Veterans Affairs, and other agencies announced that
22 they would terminate or suspend use of Anthropic’s products. Each cited only the Presidential
23 Directive. *See* AR 321 (State); AR 256 (Treasury); AR 527 (VA); AR 317 (FHFA).

24 Less than two hours after the Presidential Directive, Secretary Hegseth issued a “final”
25 decision (the “Hegseth Directive”) via social media. AR 255B. He accused Anthropic of
26 “betrayal”; declared its conduct “fundamentally incompatible with American principles”;
27 “direct[ed] the Department of War to designate Anthropic a Supply-Chain Risk to National
28 Security”; ordered that, “[e]ffective immediately, no contractor, supplier, or partner that does

1 business with the United States military may conduct any commercial activity with Anthropic”
2 (the “secondary boycott” order); and instructed that Anthropic would continue to provide its
3 technology to the Department for up to six months. AR 255B.

4 Later, the Secretary signed a determination dated March 3 (the “March Determination”)
5 that formalized his designation of Anthropic as a supply chain risk under 10 U.S.C. § 3252.
6 AR 209. But the Department did not reveal that determination to Anthropic on March 3. Instead, it
7 continued negotiations with Anthropic about the disputed contract terms. Heck Decl. ¶¶ 20-22.
8 Those negotiations extended through the morning of March 4, when Under Secretary Michael sent
9 Anthropic a revised counteroffer and wrote: “I think we are very close here.” *Id.* Ex. 1; *see also*
10 Dkt. 164 ¶ 104 (admitting that negotiations extended through the evening of March 4). That
11 evening, however, Anthropic received a letter (dated March 3) notifying it of the March
12 Determination. AR 237-38; Heck Decl. ¶ 21. The Department did not disclose the text of that
13 Determination—or any materials purporting to justify it—until March 17. *See* Dkt. 96-2.²

14 On March 26, this Court preliminarily enjoined Defendants’ implementation of the
15 Presidential Directive, the Hegseth Directive, and the March Determination (together, the
16 “Challenged Actions”). Dkt. 134 (“Op.”). The Court restored the pre-February 27 status quo,
17 while also expressly allowing Defendants to take any lawful step to halt their work with Anthropic
18 and its AI tools. Dkt. 135 at 3. In the months since the Challenged Actions, the Department and
19 other federal agencies have continued using those tools. Ramasamy Decl. ¶¶ 78-81. A senior
20 official also publicly acknowledged that the Department has used Claude in active military
21 operations in Iran. Mongan Decl. Ex. 5 at 28. And before making a new and powerful model
22 named Mythos more broadly available, Anthropic briefed senior officials across government
23 agencies on its capabilities and made it available to government partners. Ramasamy Decl. ¶¶ 80-
24 81; Mongan Decl. Ex. 3.

25
26
27
28

² The Department sent Anthropic another letter on March 4 notifying it of a separate determination
by the Secretary under 41 U.S.C. § 4713. Anthropic is currently challenging that action in a
separate proceeding in the D.C. Circuit, as required by 41 U.S.C. § 1327(b)(3).

ARGUMENT

This Court should grant summary judgment in favor of Anthropic on its constitutional claims because “there is no genuine issue as to any material fact,” and Anthropic “is entitled to judgment as a matter of law.” *Anderson v. Liberty Lobby, Inc.*, 477 U.S. 242, 250 (1986). As to Anthropic’s APA claims, this Court reviews the challenged final agency action pursuant to 5 U.S.C. § 706. *Compassion Over Killing v. FDA*, 849 F.3d 849, 854 (9th Cir. 2017). Anthropic should prevail on its APA claims whether they are viewed in light of the proffered administrative record, *see* Dkt. 156-1, or with the benefit of additional documents that the Court may consider under these circumstances, *see infra* pp. 20-21. Either way, the actions are invalid.

I. Anthropic Is Entitled To Summary Judgment On Its First Amendment Claim Because Defendants Retaliated Against Anthropic For Its Protected Expression

To prevail on its First Amendment claim, Anthropic must establish “that (1) it engaged in constitutionally protected activity; (2) the defendant’s actions would ‘chill a person of ordinary firmness’ from continuing to engage in the protected activity; and (3) the protected activity was a substantial motivating factor in the defendant’s conduct.” *Ariz. Students’ Ass’n v. Ariz. Bd. of Regents*, 824 F.3d 858, 867 (9th Cir. 2016). Once Anthropic has made that showing, “the burden shifts to the defendant official to demonstrate that even without the impetus to retaliate he would have taken the action complained of.” *Hartman v. Moore*, 547 U.S. 250, 260 (2006).

The undisputed facts here establish a “classic” case of “illegal First Amendment retaliation.” Op. 2. This Court has already concluded that Anthropic “engaged in protected activity.” *Id.* at 20. Anthropic has been “a loud and influential voice regarding the capabilities, risks, and safe uses of AI technology.” *Id.*; *see also* Kaplan Decl. ¶¶ 8-10, 13; Heck Decl. ¶¶ 5-6; Mongan Decl. Exs. 1, 4. Defendants have introduced no contrary evidence, and they have never disputed that the Challenged Actions would chill a person of ordinary firmness from engaging in that protected activity. Op. 21. And the very words used by President Trump, Secretary Hegseth, and Under Secretary Michael in announcing and justifying those actions leave no question that Anthropic’s protected activity substantially motivated them. In light of those words—and other

undisputed evidence—Anthropic has carried its burden on its retaliation claim, and Defendants cannot show they would have taken the Challenged Actions absent the retaliatory impetus.

A. The Undisputed Facts Establish A Prima Facie Case Of Retaliation

1. Anthropic Engaged In Constitutionally Protected Activity

The evidence shows that Anthropic engaged in activity protected by the First Amendment. Since its founding, Anthropic has been a leading voice in the ongoing public debate about AI safety. It has worked to advance AI research, safety, and policy, *see* Kaplan Decl. ¶¶ 9-10, 13, including by publishing essays and public statements about the risks of deploying AI without proper restrictions and safeguards, *see id.* ¶¶ 10, 13; Mongan Decl. Exs. 1, 4. Its CEO has been a particularly prominent advocate for AI safety, including through a January 2026 public essay calling for “guardrails” on the development of powerful AI to address its risks to public safety and security. Mongan Decl. Ex. 4 at 21, 22, 30-32, 44. That speech on matters of profound ““public concern”” lies at the core of what the First Amendment protects. *Snyder v. Phelps*, 562 U.S. 443, 451-52 (2011).

The undisputed facts demonstrate that Anthropic’s interactions with the Department are also protected First Amendment activity. When summoned to Washington on February 24 and pressed by the Secretary to abandon Anthropic’s longstanding restrictions against the use of Claude for mass surveillance of Americans or lethal autonomous warfare, Anthropic’s leadership explained to the government the importance of those restrictions and respectfully declined to abandon them. Heck Decl. ¶¶ 13-20. Anthropic was addressing “governmental conduct” regarding the safe use of AI—a subject in which the public has a “deep and abiding interest” because it “affects the societal interest as a whole.” *Alpha Energy Savers, Inc. v. Hansen*, 381 F.3d 917, 926-27 (9th Cir. 2004). The communications “touche[d] on fundamental questions” of “policy.” *Janus v. Am. Fed’n of State, Cty., & Mun. Emps., Council 31*, 585 U.S. 878, 913 (2018). And they implicated the ““ethical [and] moral”” principles that define Anthropic’s public-benefit mission. *Burwell v. Hobby Lobby Stores, Inc.*, 573 U.S. 682, 701 (2014). The speech was no less protected just because Anthropic expressed itself directly to the government: The Free Speech Clause protects speech on “matters of substantial public concern” during contractual bargaining with the

1 government. *See Janus*, 585 U.S. at 885-86. And the Petition Clause independently protects efforts
 2 to advocate before “all departments of the Government.” *Cal. Motor Transp. Co. v. Trucking*
 3 *Unlimited*, 404 U.S. 508, 510 (1972).

4 The First Amendment likewise protects Anthropic’s public communications about the
 5 dispute. On February 26, Anthropic’s CEO issued a statement explaining that Anthropic could not
 6 “in good conscience” abandon its two restrictions because certain uses of AI fall “outside the
 7 bounds of what today’s technology can safely and reliably do.” Mongan Decl. Ex. 1 at 1-2. That
 8 public statement on a question of pressing national importance—delivered when both sides had
 9 “publicly staked out their positions and engaged in a public debate,” Op. 20—is precisely the kind
 10 of contribution to “uninhibited, robust, and wide-open” debate that the First Amendment protects,
 11 *New York Times Co. v. Sullivan*, 376 U.S. 254, 270 (1964). Indeed, it “occupies the highest rung
 12 of the hierarchy of First Amendment values.” *Snyder*, 562 U.S. at 452.

13 **2. Defendants’ Actions Would Chill A Person Of Ordinary Firmness**

14 The Challenged Actions “easily qualify as ones which would chill a person of ordinary
 15 firmness from continuing to engage in further protected speech.” Op. 21. Defendants have never
 16 contested that issue. *Id.* Nor could they. The Presidential Directive categorically bars Anthropic
 17 from working with all federal agencies, AR 255A; the secondary boycott order prohibits any
 18 “contractor, supplier, or partner that does business with the United States military” from engaging
 19 in any “commercial activity” with Anthropic, AR 255B; and the supply chain risk designation
 20 labels Anthropic a national-security threat and prevents it from competing for certain Department
 21 contracts, including as a subcontractor, AR 209. Together, the Challenged Actions “chill public
 22 debate” (Op. 21) by sending an unmistakable warning to companies in the industry: speak out and
 23 you too may be blacklisted, boycotted, and branded a national-security risk.

24 **3. Anthropic’s Speech Substantially Motivated The Challenged Actions**

25 The record does not allow for any genuine dispute that Anthropic’s protected expression
 26 was a substantial motivating factor for the Challenged Actions. Defendants’ own words establish
 27 that they acted to punish Anthropic for its public stance on whether its models could safely be
 28 used in the two contexts at issue: Secretary Hegseth tied the supply chain risk designation and the

1 secondary boycott to Anthropic’s “sanctimonious rhetoric” and its purported attempt to “strong-
 2 arm the United States military” to accept its “Silicon Valley ideology.” AR 255B. President
 3 Trump similarly explained his Directive by describing Anthropic as a “RADICAL LEFT, WOKE
 4 COMPANY” staffed by “Leftwing nut jobs” who made a “DISASTROUS MISTAKE trying to
 5 STRONG-ARM the Department of War.” AR 255A. And Under Secretary Michael doubled down
 6 in his “Risk Analysis.” AR 213. He asserted that Anthropic’s “risk level” supposedly “escalated”
 7 when the company “began engaging in an increasingly hostile manner through the press.” AR 215.
 8 He dismissed Anthropic’s position as “brand-building” and “marketing,” characterized the dispute
 9 as a “public[] spat,” and accused Anthropic of “leveraging” the parties’ “negotiations for
 10 Anthropic’s own public relations.” AR 214-15.

11 These statements expressly take issue with Anthropic’s decision to participate in “free
 12 discussion of governmental affairs”—precisely what the First Amendment was designed to protect
 13 by “practically universal agreement.” *Mills v. Alabama*, 384 U.S. 214, 218 (1966). When direct
 14 “evidence that the defendants expressed opposition to” Anthropic’s speech appears in the same
 15 statements that announced and justified the Challenged Actions, that is powerful proof that
 16 “retaliation was a substantial or motivating factor.” *Alpha Energy Savers*, 381 F.3d at 929.

17 Other factors that the Ninth Circuit considers point decisively in the same direction. The
 18 “proximity in time between [the] expressive conduct and the allegedly retaliatory action,” *id.*,
 19 confirms that the conduct motivated the actions. President Trump and Secretary Hegseth issued
 20 their Directives fewer than 24 hours after Anthropic’s CEO publicly defended its usage
 21 restrictions and declined to abandon them. Mongan Decl. Exs. 1, 7-8. If “three to eight months is
 22 easily within a time range that can support an inference of retaliation,” *Coszalter v. City of Salem*,
 23 320 F.3d 968, 977 (9th Cir. 2003), a single day certainly is.

24 And the undisputed facts also show that Defendants’ “proffered explanations for their
 25 adverse actions were false or pretextual.” *Alpha Energy Savers*, 381 F.3d at 929. To the extent the
 26 February 27 Directives offer any rationale beyond invective and criticism of First Amendment
 27 activity, they suggest a concern that Anthropic will “veto” the “operational decisions of the United
 28 States military.” AR 255B; AR 255A. The Michael memorandum likewise asserts that Anthropic

1 possesses an “operational veto” over Claude after it is deployed in the Department’s classified
 2 systems. AR 213. But none of those documents explains why Defendants “believed Anthropic has
 3 ‘privileged access,’ ‘backdoors,’ or can otherwise ‘control’ Claude after its delivery to DoW.”
 4 Op. 14; *see also* Ramasamy Decl. ¶¶ 48-58.

5 The record directly contradicts the government’s after-the-fact rationales. Once deployed
 6 in secure Department environments, Anthropic has no access to its models, no visibility into their
 7 use, and no technical ability to alter or disable them. Ramasamy Decl. ¶¶ 48-58. It cannot push
 8 changes, introduce vulnerabilities, or interfere with operations. *Id.* It cannot veto Department
 9 uses—and it has expressly disavowed any desire to interfere with military operations, including by
 10 offering a contract term expressly disavowing any authority to interfere with ongoing operations.
 11 Heck Decl. ¶ 25; *see* Ramasamy Decl. ¶¶ 48-58. Bare assertions of theoretical risk made “without
 12 evidence in the record” cannot defeat summary judgment. *White v. Lee*, 227 F.3d 1214, 1241 (9th
 13 Cir. 2000). That is especially true where (as here) those assertions are “blatantly contradicted by
 14 the record,” such that “no reasonable jury could believe” them. *Scott v. Harris*, 550 U.S. 372, 380
 15 (2007). Far from creating a triable issue, the inaccurate and pretextual nature of Defendants’
 16 chosen explanations confirms what is already apparent: their actions were substantially motivated
 17 by a desire to punish Anthropic for its protected constitutional activities.

18 **B. Defendants Cannot Rebut The Prima Facie Case**

19 Defendants can “‘escape liability only’” by proving that they “‘would have reached the
 20 same decision . . . even in the absence of’” that speech. *Thomas v. County of Riverside*, 763 F.3d
 21 1167, 1169 (9th Cir. 2014). It is not enough for Defendants to show that they “could have” taken
 22 the Challenged Actions in response to the contracting impasse. *CarePartners, LLC v. Lashway*,
 23 545 F.3d 867, 877 (9th Cir. 2008). The question is whether, without Anthropic’s protected speech,
 24 Defendants *would have* taken these same sweeping actions anyway. Op. 23. The undisputed facts
 25 show they would not have.

26 The Presidential and Hegseth Directives identify the reason for the Challenged Actions on
 27 their face. Both link the government’s response to Anthropic’s public position, which the
 28 Directives characterize as an ideologically objectionable effort to “strong-arm” the government.

1 See AR 255A, 255B; *supra* p. 9. Where the government’s own documents “admit[] that they took
 2 the action directly in response to” protected speech, “[t]here is no genuine issue of disputed fact”
 3 that the same action would not have occurred without that speech. *Riley’s Am. Heritage Farms v.*
 4 *Elsasser*, 32 F.4th 707, 728 (9th Cir. 2022). That alone should end the inquiry.

5 And nothing in the record supports the “determination that the contracting impasse alone
 6 would have triggered such an extreme response” absent Anthropic’s protected speech. Op. 24; *see*
 7 Dkt. 96 at 23; AR 213-16. If Defendants’ actual concern had been Anthropic’s refusal to accept
 8 the Department’s preferred contract terms, they had methods to address that concern directly.
 9 Defendants could have stopped using Claude; selected another vendor for the GenAI.mil contract;
 10 or used ordinary procurement authorities to terminate, narrow, or wind down particular Anthropic
 11 contracts, *see, e.g.*, 48 C.F.R. §§ 49.101(b), 49.402-3, 52.212-4(l), 52.212-4(m), 52.243-4(a),
 12 52.249-1-52.249-6. Even absent the Directives, which lay bare Defendants’ retaliatory purpose, a
 13 dispute over contract terms would not explain a government-wide blacklist, a sweeping secondary
 14 boycott, or a first-of-its-kind supply chain risk designation.

15 Recent events underscore that Defendants took the Challenged Actions in response to
 16 Anthropic’s protected speech, not genuine national-security concerns. Even after the supply chain
 17 risk designation, a Department official testified that the agency used Claude to support military
 18 efforts in Iran. Mongan Decl. Ex. 5 at 28. Anthropic has continued to work closely with the
 19 federal government, including by proactively extending access to its new Mythos model and
 20 meeting with senior White House officials in April to discuss its responsible use—discussions the
 21 White House described as “productive and constructive.” Heck Decl. ¶ 27; Mongan Decl. Ex. 7 at
 22 1. These are not the actions of parties genuinely concerned that Anthropic poses a national-
 23 security risk.

24 Finally, *American Federation of Government Employees, AFL-CIO v. Trump*, 167 F.4th
 25 1247 (9th Cir. 2026) (*AFGE*), does not help Defendants. *See* Op. 23 n.12. The government would
 26 have issued the challenged executive order in that case even absent the asserted retaliatory motive
 27 because the order itself contained no retaliatory language, it had ““legitimate grounding in national
 28 security concerns,”” and the record showed it would have issued one way or the other. *AFGE*, 167

1 F.4th at 1256-57. This case is the opposite: The Presidential and Hegseth Directives tied the
 2 Challenged Actions to Anthropic’s speech activities. And there is no record-supported national-
 3 security rationale that would independently explain the Challenged Actions. *See infra* pp. 17-21.

4 **II. Secretary Hegseth’s Actions Violate The APA Because They Are Arbitrary,**
 5 **Capricious, And Contrary To § 3252’s Procedural And Substantive Requirements**

6 On February 27, Secretary Hegseth issued a “final” decision, “[e]ffective immediately.”
 7 AR 255B. That Directive contained two unprecedented orders: The Secretary ordered a secondary
 8 boycott, prohibiting all military contractors from conducting business with Anthropic; and he
 9 “direct[ed] the Department of War to designate Anthropic a Supply-Chain Risk to National
 10 Security.” AR 255B. Both orders violate the APA. Defendants have “concede[d]” that the
 11 secondary boycott order “was in excess of and contrary to [the Secretary’s] statutory authority.”
 12 Op. 35. And the supply chain risk designation is procedurally and substantively flawed. That is
 13 true regardless of whether the Court views the relevant agency action as the Secretary’s February
 14 27 decision or his subsequent March Determination formalizing the designation, and regardless of
 15 whether the Court reviews the action exclusively in light of the administrative record proffered by
 16 Defendants or with the benefit of additional materials it may properly consider. *See infra* pp. 20-
 17 21.

18 **A. Secretary Hegseth’s February 27 Directive Is Final And Unlawful Agency**
 19 **Action**

20 To date, Defendants have offered no defense of the actions announced in Secretary
 21 Hegseth’s February 27 Directive. They instead contend that the Directive is unreviewable because
 22 it “is *not* a final agency action subject to APA review.” Dkt. 96 at 25. As this Court previously
 23 concluded, that is wrong. Op. 35. The Directive reflects the “‘consummation’ of the agency’s
 24 decisionmaking process.” *Bennett v. Spear*, 520 U.S. 154, 178 (1997). It states unequivocally that
 25 “[t]his decision is final” and “[e]ffective immediately.” AR 255B. That language conveys a
 26 “definitive decision,” not a “tentative or interlocutory” one. *U.S. Army Corps of Eng’rs v. Hawkes*
 27 *Co.*, 578 U.S. 590, 597-98 (2016). For the same reason, the Directive imposes “legal
 28 consequences” and determines “rights” and “obligations.” *Bennett*, 520 U.S. at 178; *see* Op. 35. It

1 therefore satisfies the APA’s finality requirement.

2 Neither action announced in the February 27 Directive can survive APA review. At the
 3 preliminary injunction hearing, “Defendants conceded . . . that there is no statutory basis for” the
 4 secondary boycott order. Op. 34. As to Secretary Hegseth’s order that the Department designate
 5 Anthropic a supply chain risk, Defendants have submitted no evidence that the Secretary complied
 6 with the statutory prerequisites of § 3252 before issuing it. *See id.* The Directive itself contains
 7 scarcely any reasoning—apart from criticizing Anthropic’s “rhetoric” and “ideology,” AR 255B—
 8 and lacks any genuine support from contemporaneous record evidence. The administrative record
 9 materials proffered by Defendants almost entirely post-date the Directive. The lone exception is a
 10 due-diligence report from a third-party vendor dated February 26, 2026. AR 217-29. It assigns
 11 Anthropic an overall risk of “medium,” based in substantial part on private copyright litigation
 12 against the company and the fact that “the US Secretary of Defense threatened to label the
 13 company as a ‘supply chain risk.’” AR 219-20. That circular reasoning “does not corroborate the
 14 risks that purportedly resulted in the supply chain risk designation.” Op. 36 n.18.

15 The Secretary’s March Determination merely formalized and carried out his earlier, “final”
 16 decision “directing the Department of War to designate Anthropic a Supply-Chain Risk to
 17 National Security.” AR 255B. The Secretary did not “‘deal with the problem afresh’” on March 3,
 18 as would be required for “*new agency action*.” *DHS v. Regents of the Univ. of Cal.*, 591 U.S. 1, 21
 19 (2020). Nor did he “withdraw[] or rescind[]” the earlier Directive. Op. 35. As a result, the prior
 20 Directive is the operative agency action—and the supply chain risk designation must be evaluated
 21 based on the scant reasoning in that Directive and the record before the Secretary when he issued
 22 it on February 27. *See Regents*, 591 U.S. at 20. Because the original Directive cannot be sustained,
 23 the March Determination cannot stand.

24 **B. The March Determination Is Independently Unlawful**

25 Even crediting Defendants’ position that the March Determination is the relevant final
 26 agency action, it cannot survive APA review. The action violates clear-cut procedural
 27 requirements in 10 U.S.C. § 3252; stretches the substantive authority granted in § 3252 beyond the
 28 breaking point; and rests on arbitrary, incorrect, and contrived rationales.

1 **1. The § 3252 Designation Violates Procedures Required By Law**

2 To start, Secretary Hegseth failed to comply with mandatory procedural safeguards. *See* 5
 3 U.S.C. § 706(2)(D). Before making a supply chain risk designation, the Secretary must determine
 4 in writing that “less intrusive measures” could not “reduce” the identified risk, 10 U.S.C.
 5 § 3252(b)(2)(B), and then provide Congress “a discussion of less intrusive measures that were
 6 considered and why they were not reasonably available,” *id.* § 3252(b)(3)(B). He also must
 7 determine in writing that the designation is “necessary to protect national security by reducing
 8 supply chain risk,” *id.* § 3252(b)(2)(A), and provide Congress “a summary of the basis for the
 9 determination,” *id.* § 3252(b)(3)(B). Under the Department’s regulations, that latter determination
 10 must be based on “a risk assessment [prepared] by the Under Secretary of Defense for
 11 Intelligence.” 48 C.F.R. § 239.7304(a). Collectively, these procedural requirements help to ensure
 12 that the Secretary appropriately exercises the weighty authority granted him in § 3252.

13 The proffered administrative record shows that the Secretary violated each of those
 14 requirements here. He failed to make any reasoned determination that less intrusive measures
 15 could not reduce any risk Anthropic posed. Although the March Determination pronounces that
 16 conclusion, it does not explain it. AR 209. Nor does any other contemporaneous document in the
 17 record. And “mere lip service or verbal commendation of [the] standard” cannot satisfy the
 18 statutory requirement. *Nat. Res. Def. Council, Inc. v. Pritzker*, 828 F.3d 1125, 1135 (9th Cir.
 19 2016); *see* Op. 33-34. The Secretary also failed to provide Congress the required analysis of less
 20 intrusive measures, “as Defendants conceded at oral argument.” Op. 33; *see* P.I. Hr’g Tr. 20:15-
 21 21:1, 23:10-12. The six congressional notification letters in the record incant § 3252’s text, but do
 22 not contain the mandated “discussion.” 10 U.S.C. § 3252(b)(3)(B). Further still, the underlying
 23 risk assessment was prepared by Under Secretary for Research and Engineering Emil Michael—
 24 not the Under Secretary for Intelligence, as the Department’s regulation requires. AR 213-16, 243.
 25 Defendants seek to justify reassigning the risk assessment by pointing to a recent Department
 26 reorganization, *see* AR 244-45, but an internal reorganization is no basis to act “contrary to
 27 existing valid regulations.” *United States ex rel. Accardi v. Shaughnessy*, 347 U.S. 260, 268
 28 (1954); *see Nat’l Ass’n of Home Builders v. Norton*, 340 F.3d 835, 852 (9th Cir. 2003).

2. The § 3252 Designation Is Contrary To Law

As a substantive matter, the supply chain risk designation is incompatible with the text of § 3252. *See* 5 U.S.C. § 706(2)(A), (C). Congress defined “supply chain risk” as “the risk that an adversary may sabotage, maliciously introduce unwanted function, or otherwise subvert . . . a covered system.” 10 U.S.C. § 3252(d)(4). That text requires an “adversary” and “is directed at covert acts or hacks”—“not overt positions taken during contract negotiations.” Op. 30.

Although the March Determination does not explain how Anthropic presents the kind of risk Congress described, AR 209, the underlying risk analysis asserts that “Anthropic’s risk level escalated from a potentially manageable technical and business negotiation to an unacceptable national security threat over the course of the DoW’s contract negotiation[s] with them,” AR 215. Those negotiations, and Anthropic’s private and public statements about them, purportedly led the Department to conclude that there was a “significant risk” that Anthropic itself was an adversary that would “subvert the design and/or functionality of their product or service.” AR 210; *see also* AR 215. But Defendants ultimately assert that if Anthropic had “agreed to the Government’s term, the challenged actions would not have occurred”—meaning that Anthropic’s refusal to agree was the dispositive factor in the Secretary’s § 3252 designation. Dkt. 96 at 23. Defendants thus “appear to be taking the position that any vendor who ‘push[es] back’ on” the Department in a contract negotiation becomes an “adversary” who presents a supply chain risk. Op. 30.

As this Court has recognized, “[t]hat position is deeply troubling and inconsistent with the statutory text.” Op. 31. Under § 3252’s plain terms, the stance Anthropic took in negotiations with the Department provides no basis for a supply chain risk designation. The text reflects Congress’s choice to target the risk of malicious entities engaged in “covert acts or hacks,” *id.*, not vendors openly expressing views with which the Department disagrees. A supply chain risk must come from an “adversary”—that is, an “enemy,” “foe,” “opponent,” or “rival.” *Adversary*, Black’s Law Dictionary (12th ed. 2024). Those words naturally describe hostile foreign countries, terrorist groups, and similar actors. They do not describe a company like Anthropic, whatever one thinks of its stance on AI safety. A contractor does not give rise to a national-security risk because it disagrees with the Department about how a technology can be used safely. That is particularly so

1 where the contractor has consistently supported American national security agencies; has been
2 entrusted with access to Top Secret facilities and classified information; and continues to provide
3 the Department with technological tools that help to protect the Nation. *See supra* pp. 5-6.

4 Other statutory terms reinforce that conclusion. Section 3252 does not address all risks
5 from an “adversary” but only the risks of “sabotage,” “subver[sion],” and the “malicious[]
6 introduc[tion of] unwanted function.” 10 U.S.C. § 3252(d)(4). Each of those actions is typically
7 taken or planned covertly. *See, e.g., Sabotage*, Oxford English Dictionary (2d ed. 1989) (“[T]o
8 ruin, destroy, or disable deliberately and maliciously (frequently by indirect means).”); *Subvert*,
9 Oxford English Dictionary (2d ed. 1989) (“[T]o attempt to achieve, esp. by covert action, the
10 weakening or removal of (a government, political regime, etc.).”). The statute further limits the
11 relevant covert actions to those taken “so as to surveil, deny, disrupt, or otherwise degrade”
12 covered national security systems. 10 U.S.C. § 3252(d)(4). That language connotes purposeful
13 misconduct. *See So As*, Oxford English Dictionary (2d ed. 1989) (“Expressing result,
14 consequence, or purpose: in such a way that; so that, in order that.”). But Defendants designated
15 Anthropic over its *public* disagreement about contract language. Dkt. 96 at 23. A company that
16 openly announces its position in negotiations and publicly defends that position is not
17 “sabotag[ing],” “subvert[ing],” or “maliciously introduc[ing]” anything. Those words describe the
18 opposite of the transparent and public position taken by Anthropic here. *See Op.* 30.

19 “The legislative history confirms that Section 3252’s purpose is to address covert acts.”
20 *Op.* 31. Congress enacted § 3252’s predecessor provision in response to a cyber incident
21 perpetrated by a foreign intelligence agency. That episode sparked concerns about how “the
22 globalization of the information technology industry” could increase the risk that “systems and
23 networks critical to DOD could be exploited through the introduction of counterfeit or malicious
24 code and other defects introduced by suppliers of systems or components.” S. Rep. No. 111-201,
25 at 162 (2010). Those concerns bear no relation to the Department’s dispute with Anthropic—an
26 American company that actively supports the Department’s mission and that the Department has
27 never accused of an “attack,” *id.*, or anything of the sort.

28 “[C]onsidered and consistent Executive Branch practice” resolves any doubt. *Kennedy v.*

1 *Braidwood Mgmt., Inc.*, 606 U.S. 748, 783 (2025). Shortly after Congress enacted the provision
 2 that became § 3252, the Department issued instructions defining the relevant risk as “sabotage or
 3 subversion” by “foreign intelligence, terrorists, or other hostile elements.” Dep’t of Def.,
 4 Instruction No. 5200.44 (Nov. 5, 2012), perma.cc/TW3D-735M. In keeping with that
 5 understanding, the Department has never before designated an entity a supply chain risk due to a
 6 contract dispute.

7 **3. Defendants’ Rationales Are Arbitrary And Capricious**

8 Secretary Hegseth’s proffered explanations are also arbitrary and capricious, providing
 9 another basis for this Court to hold unlawful and set aside the § 3252 Designation. 5 U.S.C.
 10 § 706(2)(A). The March Determination itself is wholly conclusory. AR 209. The only
 11 contemporaneous rationales appear in the memorandum from Under Secretary Michael. AR 213-
 12 16. But its justifications provide nothing close to the reasoned explanation the APA requires.

13 The overarching theme of the memorandum is that Anthropic presents a supply chain risk
 14 because its “unwillingness to agree” to the Department’s preferred contract term amounted to an
 15 attempt “to grant itself” a “veto” over Department operations. AR 213. The memorandum asserts
 16 that Anthropic’s preferred terms create a risk that it will ““deny, disrupt, or otherwise degrade””
 17 covered Department systems by “diminish[ing] functionality and limit[ing] DoW’s warfighting
 18 capabilities.” *Id.* A post-decision Michael declaration reiterates the point, stating that Anthropic’s
 19 two usage restrictions amount to “an operational veto” that “confirms that Anthropic did want to
 20 intervene in and assert control over DoW operations.” AR 252; *see* P.I. Hr’g Tr. 41:18-23.

21 Those assertions are unfounded. Contract terms the Department was free to reject are not
 22 intrusions into the “operational decision chain.” AR 215. And there is nothing reasoned or
 23 reasonable about the contention that, by refusing to abandon two narrow use restrictions,
 24 Anthropic demonstrated a propensity to interfere with Department operations. A company’s
 25 forthright insistence on usage restrictions in contract negotiations (conducted to establish ground
 26 rules *before* models are deployed) does not provide evidence that the company would sabotage the
 27 Department *after* its models are deployed. A saboteur would not care about the particulars of
 28 contract terms. And the notion that the two use restrictions maintained by Anthropic signal a risk

1 of sabotage is especially irrational because the Department insists that it has no intention of using
2 Anthropic’s models for those uses. *See* Dkt. 96 at 14. Whatever deference is due to the
3 Department’s national-security judgments cannot provide cover for treating good-faith contractual
4 disagreements as evidence of threatened sabotage or subversion.

5 In the related D.C. Circuit litigation, Secretary Hegseth emphasized a somewhat different
6 rationale in Michael’s memorandum: that Anthropic is a supply chain risk because AI technology
7 is “opaque” and “vulnerable to manipulation,” and the Department “cannot trust Anthropic to
8 ensure the integrity of its models.” AR 214; *see, e.g.,* Br. for Respondents 26-27, *Anthropic v.*
9 *Dep’t of War*, No. 26-1049 (D.C. Cir. May 6, 2026). But that rationale is no better than the first.
10 The opacity of AI technology is “a common concern with all” large language models, AR 245,
11 and the possibility of vendor manipulation is a risk with “technology more broadly,” Op. 32.

12 That leaves only the “trust” rationale, which similarly fails. Michael rested his assertion
13 that Anthropic “cannot be trusted” on Anthropic’s purported conduct during negotiations. AR 214.
14 He viewed Anthropic’s negotiating posture as designed “to benefit [Anthropic’s] public
15 perception,” and objected to Anthropic’s engagement with “the press” and to a single question
16 from an Anthropic employee about the Department’s use of Claude. AR 213-15. In his view, that
17 conduct created “material doubts as to whether [Anthropic] would cause [its] software to stop
18 working or cause some other disastrous action.” AR 213.

19 Even crediting that characterization of Anthropic’s conduct, there is no “rational
20 connection” between the conduct and Michael’s conclusion. *Burlington Truck Lines Inc. v. United*
21 *States*, 371 U.S. 156, 168 (1962). Every company wants to burnish its public image; that does not
22 make it a threat of sabotage or covert attack. *See* AR 214. Nor is it reasonable to infer that
23 Anthropic might undermine future military operations because the company raised a question
24 about a prior operation or criticized the Department in the press. Engaging in public disagreements
25 and frank conversations with the Department is a most unlikely strategy for any entity intent on
26 “sabotag[ing]” or “subvert[ing]” it. 10 U.S.C. § 3252(d)(4). While Defendants might have lost
27 “trust” in Anthropic as a contracting partner, that hardly establishes that Anthropic would secretly
28 and maliciously interfere with the Department’s information systems.

1 Indeed, the proffered administrative record reveals that the Department’s explanations for
2 its actions are pretextual. Almost all of the records produced in support of the designation were
3 created *after* Secretary Hegseth ordered the Department to designate Anthropic a supply chain risk
4 on February 27. *See* Op. 36; *supra* p. 13. That timeline “raises an inference” that the record “was
5 generated after the fact to justify” a “foreordained conclusion.” Op. 36.

6 The Department’s own contradictory statements further demonstrate that this is a
7 designation in search of a justification. According to the Michael memorandum, “Anthropic’s risk
8 level escalated” during its dealings with the Department to the point when, “during the final weeks
9 of negotiations,” the company became a “fully mature supply chain risk.” AR 215. Yet
10 Defendants also insist that if Anthropic had just accepted the Department’s contract term on
11 February 27, “the challenged actions would not have occurred.” Dkt. 96 at 23. It is inconceivable
12 that Defendants would allow a company that supposedly presented a “fully mature” threat of
13 “model poisoning, insider threat risk, data exfiltration, and denial of service” (AR 215) to continue
14 providing services merely because it agreed to a contract term. *See also* AR 255B (February 27
15 Hegseth Directive allowing continued use of Anthropic’s services for up to six months).
16 “Altogether, the evidence tells a story that does not match the explanation the Secretary gave for
17 his decision.” *Dep’t of Com. v. New York*, 588 U.S. 752, 784 (2019).

18 The “belated justifications” offered by Under Secretary Michael in his declarations filed in
19 this litigation on March 17 and March 24 do not salvage the Secretary’s action. *Regents*, 591 U.S.
20 at 23. Agency action cannot be upheld based on “*post hoc* rationalizations.” *Am. Textile Mfrs.*
21 *Inst., Inc. v. Donovan*, 452 U.S. 490, 539 (1981). In any event, the post hoc rationalizations here
22 would not change the analysis. The concern about Anthropic’s employment of “foreign nationals”
23 applies to “other major U.S. AI labs,” as Michael concedes. AR 245. The fact that “adversarial
24 nation states” have stolen Anthropic’s technology (AR 246) makes Anthropic a victim, not an
25 adversary. And when the CDC used a commercial version of Claude to conduct public health
26 research and encountered difficulties as a result of the commercial Usage Policy, it is undisputed
27 that Anthropic promptly assisted the CDC as soon as it learned of the situation. AR 246. Far from
28 suggesting sabotage or adversarial intent, that shows a company responding to government needs

1 and working collaboratively with federal partners.

2 **C. Additional Materials Properly Considered Here Reinforce The APA Violations**

3 As addressed in a concurrently filed motion, in evaluating the APA claim, this Court
 4 should consider two additional documents beyond the materials proffered by Defendants. The
 5 first—an email from Under Secretary Michael to Anthropic—is part of the “whole record.” 5
 6 U.S.C. § 706; *see Thompson v. U.S. Dep’t of Lab.*, 885 F.2d 551, 555 (9th Cir. 1989). The
 7 second—the June 10 Declaration of Thiyagu Ramasamy (“the Ramasamy Declaration”)—may be
 8 considered “to identify and plug holes in the administrative record.” *Lands Council v. Powell*, 395
 9 F.3d 1019, 1030 (9th Cir. 2005). That additional evidence confirms the APA violations.

10 Under Secretary Michael’s correspondence with Anthropic confirms that the asserted
 11 concerns about Anthropic presenting a supply chain risk are “contrived.” *Dep’t of Com.*, 588 U.S.
 12 at 784. On March 4—two days after Under Secretary Michael prepared the risk assessment and
 13 one day after the Secretary’s Determination—Michael sent Anthropic a revised counteroffer. Dkt.
 14 114-1 at 2. He wrote: “I think we are very close here,” and “I hope this work[s].” *Id.* If the
 15 Department genuinely believed that Anthropic was likely to maliciously disrupt military
 16 operations, such continued and enthusiastic negotiations would have been unthinkable.

17 The additional materials also fatally undermine the Department’s contentions that
 18 Anthropic demanded an “operational veto,” AR 213, and that the company might “alter the
 19 behavior” of its models “in the middle of ongoing warfighting operations.” AR 215. As the
 20 Ramasamy Declaration explains, Anthropic lacks the technical capability to take any such actions.
 21 It “has no ability to access, alter, or shut down” a model that has been “deployed in DoW
 22 environments.” Ramasamy Decl. ¶ 48. It maintains no “back door or remote kill switch” in
 23 Claude. *Id.* And “Anthropic personnel cannot log into a Department system to modify or disable
 24 the models during an operation.” *Id.* Before a new or updated model is deployed, moreover,
 25 “Department security reviewers inspect and approve it to ensure it is safe and appropriate.” *Id.*
 26 ¶ 53. This evidence—all of which Anthropic would have submitted to the Department before the
 27 § 3252 designation, if given the opportunity to do so—shows that Anthropic could not exercise an
 28 “operational veto” even if it wanted to. AR 213.

III. Defendants Violated The Due Process Clause By Blacklisting Anthropic Without Notice Or An Opportunity To Be Heard

The undisputed facts also establish a violation of Anthropic's due process rights. To begin with, they show that Anthropic has "been deprived of a protected interest in 'property' or 'liberty.'" *Am. Mfrs. Mut. Ins. Co. v. Sullivan*, 526 U.S. 40, 59 (1999). Defendants previously argued (Dkt. 96 at 30) that the due process claim "fails at this first step." That argument can be rejected here because it relies on no disputed facts, only a mistaken understanding of the law.

The Challenged Actions deprived Anthropic of three protected interests. *First*, the Presidential Directive abridged Anthropic's liberty interest in "follow[ing] [its] chosen profession free from unreasonable governmental interference," *Greene v. McElroy*, 360 U.S. 474, 492 (1959), by "debarring" it "from government contract[ing]." *Trifax Corp. v. Dist. of Columbia*, 314 F.3d 641, 643 (D.C. Cir. 2003); *see* Op. 25-26. The Presidential Directive ordered federal agencies to "IMMEDIATELY CEASE all use of Anthropic's technology" and "not do business with [it] again." AR 255A. And Defendants' implementation of that ban "automatically exclude[s]" Anthropic from a "category of future [agency] contracts," likewise triggering due process protections. *Kartseva v. Dep't of State*, 37 F.3d 1524, 1528 (D.C. Cir. 1994). *Second*, the Challenged Actions deprived Anthropic of "a cognizable liberty interest," *Fikre v. FBI*, 35 F.4th 762, 776 (9th Cir. 2022), by stigmatizing Anthropic as "a Supply-Chain Risk to National Security," AR 255B, and a company that put "AMERICAN LIVES at risk," AR 255A, while stripping it of federal contracts and any opportunity to compete for contracts in the future. *See Ulrich v. City & Cty. of San Francisco*, 308 F.3d 968, 983 (9th Cir. 2002); Op. 26-27. *Third*, the secondary boycott order prohibited Anthropic from "conduct[ing] any commercial activity" with military contractors, AR 255B, implicating "significant" property interests, *Al Haramain Islamic Found. v. Dep't of Treasury*, 686 F.3d 965, 979-80 (9th Cir. 2012); *see also Ralls Corp. v. CFIUS*, 758 F.3d 296, 315-16 (D.C. Cir. 2014).

And the meager post-deprivation procedures that Defendants offered Anthropic here do not "comport with due process." *Sullivan*, 526 U.S. at 59. "The fundamental requirements of procedural due process are notice and an opportunity to be heard *before* the government may

1 deprive a person of a protected liberty or property interest.” *Conner v. City of Santa Ana*, 897 F.2d
 2 1487, 1492 (9th Cir. 1990) (emphasis added) (citing *Mathews v. Eldridge*, 424 U.S. 319, 333
 3 (1976)). The record here establishes a “complete lack of prior notice or process” as to each of the
 4 Challenged Actions. Op. 27; *see also* Heck Decl. ¶¶ 22, 24.

5 The government told Anthropic that it wanted to alter the parties’ contract terms, *see* Heck
 6 Decl. ¶ 11, but gave Anthropic no notice or opportunity to object before publicly barring the
 7 company from all federal government work and blacklisting it with all military contractors, *see*
 8 AR 255A, 255B. The Department also did not provide Anthropic with the factual basis for the
 9 supply-chain-risk designation before imposing it. The document containing the purported
 10 rationales was not disclosed to Anthropic until after this litigation began, *see* AR 213-16, and that
 11 analysis relied on accusations that Anthropic had no prior opportunity to contest, AR 214-15. Nor
 12 did Secretary Hegseth’s statements during negotiations provide the requisite notice. At the same
 13 time he suggested that Anthropic could be designated a supply chain risk, he also suggested that
 14 the Department could invoke the Defense Production Act to force Anthropic to provide its
 15 technology on the Department’s preferred terms. Heck Decl. ¶ 17. Those conflicting positions did
 16 not give Anthropic meaningful notice of the basis for the later designation, much less a
 17 meaningful opportunity to respond before Defendants acted. *See* Op. 27 (opining that “[t]his
 18 contradiction underscores the complete lack of prior notice or process.”).

19 Anthropic also did not receive pre-deprivation notice or process from the agencies that
 20 carried out the Presidential Directive. Within hours of that Directive, GSA directed Anthropic’s
 21 reseller to submit a deletion modification, AR 403, and removed Anthropic from USAi.gov and
 22 the Multiple Award Schedule, AR 394-95, 408. HHS disabled enterprise access to Claude that
 23 evening. AR 335, 336. The State Department initiated removal of the Anthropic model from its
 24 internal generative AI tool by the close of business. AR 321. None of these actions was preceded
 25 by notice to Anthropic, an opportunity to respond, or any agency-specific risk finding.

26 Defendants contend that they provided constitutionally sufficient post-deprivation process.
 27 *See* Dkt. 96 at 31. But they did not offer *any* post-deprivation process for the Presidential
 28 Directive or Secretary Hegseth’s secondary boycott order. And the 30-day reconsideration process

1 the Department offered on March 4 with respect to the § 3252 designation (*see* AR 238) is
 2 constitutionally inadequate. The constitutional baseline is that “[p]re-deprivation notice typically
 3 is required absent ‘a strong justification’ from the government.” *Garza v. Woods*, 150 F.4th 1118,
 4 1130 (9th Cir. 2025); *see generally Cleveland Bd. of Educ. v. Loudermill*, 470 U.S. 532, 542
 5 (1985). The record here establishes no such justification. *See* Op. 28-29; AR 209-16.

6 In Defendants’ telling, the “national security” setting warranted only post-deprivation
 7 process. Dkt. 96 at 31. But “the foreign-policy/national-security nature” of an action does not
 8 automatically “support[] the constitutional adequacy of a post-deprivation remedy.” *Nat’l Council*
 9 *of Resistance of Iran v. Dep’t of State*, 251 F.3d 192, 207 (D.C. Cir. 2001) (*NCRI*). Indeed, the
 10 D.C. Circuit has repeatedly held that due process typically requires advance notice as to the
 11 government’s intent to designate an entity a foreign terrorist organization. *See id.* at 207-08;
 12 *People’s Mojahedin Org. of Iran v. Dep’t of State*, 613 F.3d 220, 227-28 (D.C. Cir. 2010) (per
 13 curiam). Even in the national-security context, therefore, Defendants must “make a showing of
 14 particularized need” to dispense with pre-deprivation process. *NCRI*, 251 F.3d at 208. Situations
 15 presenting such a need are “‘extraordinary’” and “truly unusual.” *Fuentes v. Shevin*, 407 U.S. 67,
 16 90 (1972); *see, e.g., Al Haramain*, 686 F.3d at 985 (substantial risk of “asset flight”).

17 Here, Defendants have “present[ed] no evidence of exigency whatsoever,” Op. 28, let
 18 alone a “particularized need” sufficient to justify depriving Anthropic of process before the
 19 Challenged Actions, *NCRI*, 251 F.3d at 208. To be sure, they have vaguely suggested a need for
 20 “quick action.” Dkt. 96 at 31. But “nothing in the record indicates a heightened supply chain risk
 21 to national security systems on February 27 or March 3.” Op. 28. To the contrary, the Department
 22 had been using Claude subject to the restrictions to which it now objects. Ramasamy Decl. ¶¶ 8,
 23 25-27, 65; Kaplan Decl. ¶ 29. The Department’s own conduct refutes any urgency: it allowed
 24 Anthropic six months to wind down, AR 242, and it continues to rely on Anthropic’s tools for
 25 national-security functions to this day, Heck Decl. ¶ 27; Ramasamy Decl. ¶¶ 80-81; *see* Mongan
 26 Decl. Exs. 2-3, 6.

27 The record also “show[s] that meaningful pre-deprivation notice of DoW’s concerns and
 28 an opportunity to respond to those concerns were necessary to avoid the risk of erroneous

deprivation.” Op. 27-28; *see Mathews*, 424 U.S. at 335. The central rationale underlying the § 3252 designation is the assertion that Anthropic would alter or shut down its models to compromise Department operations. Op. 28; *see* AR 213, 215. But the unrebutted evidence now “shows that Anthropic is not technologically capable of taking such unilateral actions.” Op. 28; *see supra* pp. 10, 21; Ramasamy Decl. ¶¶ 48-58. Given the opportunity, Anthropic would have corrected the Department’s misunderstanding before the Department took action. Instead, the Department’s “factual errors [went] uncorrected” because its failure to afford constitutionally required process deprived the Secretary of “easy, ready, and persuasive explanations.” *Al Haramain*, 686 F.3d at 982.

IV. The Presidential Directive Violates The Separation Of Powers

The Presidential Directive is an immediate, permanent, government-wide debarment of Anthropic, issued by executive fiat. Presidential power must “stem either from an act of Congress or from the Constitution itself.” *Youngstown Sheet & Tube Co. v. Sawyer*, 343 U.S. 579, 585 (1952); *see also State v. Su*, 121 F.4th 1, 7-13 (9th Cir. 2024). The Directive does not. It invokes no statutory authority and has no constitutional basis. Congress—not the President—sets the rules for exclusion from federal contracting, including which contractors may be excluded, under what circumstances, and through what process. *See, e.g.*, 10 U.S.C. §§ 3203(a)(1), 3204(a); 41 U.S.C. §§ 3303(a)(1), 3304(a); *see also* 48 C.F.R. §§ 9.406-2, 9.407-2. And even the Executive has recognized that exclusion of a contractor is not a tool “for purposes of punishment.” 48 C.F.R. § 9.402(b). By defying Congress’s chosen scheme for exclusion from federal contracting, the President acted at the “lowest ebb” of presidential authority, and the Directive may stand only if it is supported by an independent constitutional power. *Youngstown*, 343 U.S. at 637 (Jackson, J., concurring). It is not.

V. Implementation Of The Presidential Directive Violated 5 U.S.C. § 558

Finally, the actions taken by defendant agencies to implement the Presidential Directive violated the APA. Agencies may not impose sanctions or issue orders “except within jurisdiction delegated to the agency and as authorized by law.” 5 U.S.C. § 558(b); *see Am. Bus Ass’n v. Slater*, 231 F.3d 1, 7 (D.C. Cir. 2000). No statute authorizes federal agencies to abruptly impose blanket

orders and sanctions that terminate all Executive Branch use of Anthropic tools. *See, e.g.*, AR 259-66; AR 319-32; AR 403-10. Those actions violate § 558(b) and must be set aside.

CONCLUSION

Anthropic's motion for summary judgment should be granted.

Respectfully submitted,

Date: June 10, 2026

/s/ Michael J. Mongan

KELLY P. DUNBAR (*pro hac vice*)
 kelly.dunbar@wilmerhale.com
 JOSHUA A. GELTZER (*pro hac vice*)
 joshua.geltzer@wilmerhale.com
 KEVIN M. LAMB (*pro hac vice*)
 kevin.lamb@wilmerhale.com
 SUSAN HENNESSEY (*pro hac vice*)
 susan.hennessey@wilmerhale.com
 LAUREN MOXLEY BEATTY (SBN 308333)
 lauren.beatty@wilmerhale.com
 MATTHEW E. MORRIS (*pro hac vice*)
 matt.morris@wilmerhale.com
 BARDIA VASEGHI (*pro hac vice*)
 bardia.vaseghi@wilmerhale.com
 DAKOTA C. FOSTER (*pro hac vice*)
 dakota.foster@wilmerhale.com
 WILMER CUTLER PICKERING
 HALE AND DORR LLP
 2100 Pennsylvania Avenue NW
 Washington, DC 20037
 Telephone: (202) 663-6000

MICHAEL J. MONGAN (SBN 250374)
 michael.mongan@wilmerhale.com
 WILMER CUTLER PICKERING
 HALE AND DORR LLP
 50 California Street, Suite 3600
 San Francisco, CA 94111
 Telephone: (628) 235-1000

EMILY BARNET (*pro hac vice*)
 emily.barnet@wilmerhale.com
 WILMER CUTLER PICKERING
 HALE AND DORR LLP
 7 World Trade Center
 250 Greenwich St
 New York, NY 10007
 Telephone: (212) 230-8800

Attorneys for Plaintiff Anthropic PBC

CERTIFICATE OF SERVICE

I hereby certify that on this 10th day of June, 2026, I electronically filed the foregoing document using the CM/ECF System, that all participants in the case are registered CM/ECF users, and that service will be accomplished by the CM/ECF system.

By: /s/ Michael J. Mongan
Michael J. Mongan

MICHAEL J. MONGAN (SBN 250374)
michael.mongan@wilmerhale.com
WILMER CUTLER PICKERING
HALE AND DORR LLP
50 California Street, Suite 3600
San Francisco, CA 94111
Telephone: (628) 235-1000

EMILY BARNET (*pro hac vice*)
emily.barnet@wilmerhale.com
WILMER CUTLER PICKERING
HALE AND DORR LLP
7 World Trade Center
250 Greenwich Street
New York, NY 10007
Telephone: (212) 230-8800

Attorneys for Plaintiff Anthropic PBC

KELLY P. DUNBAR (*pro hac vice*)
kelly.dunbar@wilmerhale.com
JOSHUA A. GELTZER (*pro hac vice*)
joshua.geltzer@wilmerhale.com
KEVIN M. LAMB (*pro hac vice*)
kevin.lamb@wilmerhale.com
SUSAN HENNESSEY (*pro hac vice*)
susan.hennessey@wilmerhale.com
LAUREN MOXLEY BEATTY (SBN 308333)
lauren.beatty@wilmerhale.com
MATTHEW E. MORRIS (*pro hac vice*)
matt.morris@wilmerhale.com
BARDIA VASEGHI (*pro hac vice*)
bardia.vaseghi@wilmerhale.com
DAKOTA C. FOSTER (*pro hac vice*)
dakota.foster@wilmerhale.com
WILMER CUTLER PICKERING
HALE AND DORR LLP
2100 Pennsylvania Avenue NW
Washington, DC 20037
Telephone: (202) 663-6000

UNITED STATES DISTRICT COURT
NORTHERN DISTRICT OF CALIFORNIA

ANTHROPIC PBC,

Plaintiff,

v.

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

Defendants.

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

DECLARATION OF JARED KAPLAN

1 I, Jared Kaplan, pursuant to 28 U.S.C. § 1746, declare as follows:

2 **Personal Background**

3 1. I am one of the co-founders of Anthropic PBC (“Anthropic”), an artificial
4 intelligence (“AI”) company based in San Francisco, California.

5 2. Since 2023, my title has been Chief Science Officer. In that role, I oversee the
6 company’s research in model development and safety, which is at the core of Anthropic’s mission.
7 Under model development, my responsibilities include overseeing the fine-tuning and
8 reinforcement-learning-driven capabilities that shape each new generation of models. I also
9 oversee teams working on interpretability, which is the study of how large language models
10 (“LLMs”) work by observing their internal operations. As part of my safety portfolio, I supervise
11 teams working on “alignment,” a term that broadly refers to efforts to make AI systems’ goals,
12 behaviors, and outputs reliably follow human values and intentions. I also oversee the safeguards
13 implemented in and around deployed models.

14 3. As of 2024, I have also served as Anthropic’s Responsible Scaling Officer. In that
15 role, I am responsible for overseeing the implementation of Anthropic’s Responsible Scaling
16 Policy, which is a series of technical and organizational protocols that aim to manage the risks
17 associated with developing increasingly capable AI systems.

18 4. Before co-founding Anthropic, I was a consultant at OpenAI and contributed to the
19 development and analysis of LLM research. During that time, I was involved in the research and
20 development efforts to introduce some of OpenAI’s early LLM models, such as GPT-3 and
21 Codex.

22 5. I began my career as a theoretical physicist, with a focus on effective field theory,
23 particle physics, cosmology, scattering amplitudes, and the conformal field theory bootstrap. Since
24 2012 and continuing to today, I have been a professor in the Department of Physics and
25 Astronomy at Johns Hopkins University.

26 6. During a sabbatical from my work in theoretical physics, I began collaborating with
27 computer scientists to research machine learning and the development of LLMs, which are text-
28 based AI systems trained on extremely large datasets to develop a functional understanding of

1 language and generate new text. Since then, I have taught courses and published over 60 scholarly
2 articles on a mixture of theoretical physics, machine learning, and LLMs.

3 7. I have extensive personal knowledge regarding Anthropic's core research and safety
4 objectives, including how Anthropic's AI models are developed and trained, their technical
5 capabilities and risks, our approach to safety research and safe deployment, and how our mission
6 informs our work across all of these domains. In my role as Chief Science Officer, I have personal
7 knowledge of the contents of this declaration, or have knowledge of the matters based on my
8 review of information and records gathered by Anthropic personnel, and could testify thereto.

9 **Anthropic's Background As An AI Safety Company**

10 8. We founded Anthropic because we anticipated how powerful AI could become and
11 believed it could reshape society in profound ways. From our own experience working with LLMs
12 and scaling laws, it became clear that AI capabilities were advancing rapidly and could soon rival
13 or surpass human performance across many areas. At the same time, we did not yet know how to
14 reliably make these systems both helpful and safe, and we anticipated that speed, competition, and
15 social disruption could push people to deploy AI before its capabilities and risks were understood
16 and sufficiently mitigated. This is why we built Anthropic—to put safety at the center as AI
17 progress accelerates, to study these questions on the most advanced models where they matter
18 most, and to build an organization that could turn careful, empirical safety research into real-world
19 practice. That conviction is embedded in Anthropic's very structure as a public benefit
20 corporation.

21 9. Our mission to build safe, beneficial AI is the foundation of everything we do—
22 from model development to safety science to policy engagement.

23 10. Anthropic began as a research-first company and, for its first two years, focused
24 exclusively on foundational AI research, the science of AI safety, and AI policy work. We
25 regularly publish pioneering research on alignment, interpretability, and the societal impacts of
26 LLMs.

27 11. Anthropic aims to develop models that are safe, ethical, and helpful. Our safety
28 work is grounded in empiricism, rigorous research, and humility. Because we have invested

1 deeply in understanding the capabilities and limitations of our systems, we have unique insights
2 into what guardrails are necessary for safe deployment.

3 12. While we initially developed AI models to support safety research, we expanded the
4 company's focus to include commercial deployments which began in early 2023. We believe this
5 matters because a safe AI system that is not used cannot fully demonstrate the benefits of
6 responsibly developed frontier AI. By showing that AI can be both safe and commercially
7 successful, we aim to pull the broader industry toward higher safety standards—what we call the
8 race to the top.

9 13. We also engage in public advocacy for transparency and safety in AI development
10 and have supported state and federal legislation advancing those goals.

11 **Anthropic's LLM Claude And The Role Of Guardrails**

12 14. Anthropic's signature model is a general-purpose LLM called Claude. We make
13 Claude available to individual users, small businesses, and large organizations through a variety of
14 offerings. We continually develop and release increasingly capable versions of Claude.

15 15. LLMs like Claude are algorithmic systems trained on massive datasets to identify
16 patterns and associations in language and to generate outputs and take actions that resemble
17 human responses and actions. Through training, models acquire predictive power and the
18 transformative ability to take a range of actions in a fraction of the time it would take humans to
19 perform them.

20 16. Claude is a versatile technology, much like an actual human mind. When paired
21 with a chatbot interface, Claude is capable of interpreting and responding to a wide range of user
22 inputs, or "prompts," in an intelligent, human-like manner. In this medium, Claude can analyze
23 and summarize large volumes of text, write and edit content, generate and debug source code, and
24 reason through complex and multi-step problems. Claude can also be given access to tools so that
25 it can behave "agentially," meaning it can not only respond to users' prompts but actually take
26 actions on their behalf. Simple actions could include sending emails, deploying code, and
27 navigating the Internet. Claude can also power more sophisticated agentic work. For instance,
28 someone planning a vacation can direct Claude to compare flight and hotel options against

1 specified constraints, assemble a day-by-day itinerary, make reservations, send confirmation
2 emails and calendar invites, and produce a summary of the plan and costs. With some
3 configurations, Claude can even act autonomously, executing tasks without requiring ongoing user
4 direction. Using AI systems to power agents is understandably of particular interest to many users,
5 including certain government users, even though agentic usage poses heightened risks relative to
6 the traditional chatbot form.

7 17. Because Claude is a dual-use technology, the risks that it poses depend on the
8 specific context in which it is used. Many tools are dual use: For example, a chef's knife is a
9 useful device in the kitchen and a dangerous weapon in a violent conflict. When it comes to AI,
10 the same capabilities that drive medical breakthroughs, accelerate scientific research, and enhance
11 human creativity can, in other contexts, also enable dangerous actors to develop new weapons or
12 automate complex, malicious activities. The dual-use nature of AI makes it difficult to restrict
13 "dangerous capabilities" while promoting "beneficial ones"—they are often the same capability
14 applied to different ends.

15 18. Although AI systems pose the potential for tremendous benefits, they also create
16 novel risks. Among other concerns, LLMs can produce responses that diverge from the goals of
17 the people that trained them or reflect skewed or mistaken judgments embedded in their training
18 data. To address the novel risks, we have pursued a multilayered approach to safety, implementing
19 safety mitigations at the model layer, the safeguards layer, and the policy layer.

20 19. At the model layer, we seek to embed safety considerations directly into the model
21 itself through a variety of training techniques. A central focus of our research is on solving the
22 challenge of alignment to make AI systems reliably follow human values and intentions. One of
23 our key techniques is Constitutional AI, which trains models to evaluate and revise their own
24 outputs against a set of normative principles, like balancing helpfulness against harm avoidance,
25 and respecting values such as individual privacy and political freedom. In addition to
26 Constitutional AI, we use reinforcement learning from human feedback (RLHF) to reduce the
27 likelihood of harmful outputs. RLHF is a training technique in which human reviewers rank pairs
28

1 of model outputs; those preferences are used to train the model to generate responses that better
2 align with human judgments.

3 20. The safeguards layer consists of technical measures that stack on top of the model
4 itself. As appropriate, our tools may include classifiers and probes that detect harmful activity in
5 real time, targeted interventions that reduce the likelihood of harmful outputs, as well as
6 monitoring systems that help us identify when our systems are being misused at scale. We
7 continually calibrate which measures are appropriate based on the circumstances of the
8 deployment, the type of harm we are trying to prevent, and other factors.

9 21. At the uppermost layer, our Usage Policy defines how Claude is permitted to be
10 used. Claude is only available subject to Terms of Service that incorporate its Usage Policy. At a
11 high level, our policy informs the development of our technical safety measures, provides users
12 with clarity on the scope of permissible usage, and steers them away from using our models in
13 risky ways, including ways we, as Claude's creator, understand that it has not been developed for
14 and/or is not ready for.

15 22. For commercial and civilian users, the Usage Policy reflects our judgment—based
16 on technical expertise, our experience at the frontier of AI development, and our values as a
17 company—on how to strike an optimal balance between enabling beneficial uses of AI while
18 mitigating potential harms. The Usage Policy generally prohibits uses that pose unacceptable
19 risks, including surveillance, compromising computer systems or networks, and designing
20 weapons or other systems to cause harm or loss of human life. The Usage Policy is an agreement
21 we enter with our users so that we both understand how Claude should and should not be used.
22 The Usage Policy is critical because technical safeguards alone cannot prevent all dangerous uses:
23 they do not necessarily have access to the full context that determines whether a given request
24 falls within or outside the lines set by the Usage Policy, and in certain environments, such as
25 classified settings, there may be very limited visibility into how the systems are being used. As a
26 result, the Usage Policy is a critical mechanism for clearly articulating safety boundaries to users
27 and serves as an important last line of defense.

23. Severing Claude from the usage limitations we have determined are essential would erode the very purpose for which our company was founded and contradict our deeply held values.

Anthropic's Commitment To Supporting National Security Engagements

While Maintaining Critical Safeguards

24. Any assertion that Anthropic is aligned with, or poses risks of subversion from, “adversaries” of the United States could not be further from the truth. We are committed to defending the United States and defeating our authoritarian adversaries. For example, we have consistently taken steps to *prevent* our models from being used by U.S. adversaries and to prioritize U.S. national security over narrow commercial self-interest. As examples, Anthropic has gone to significant lengths to prevent the use of its technology by entities linked to the Chinese Communist Party, has shut down attempts to abuse Claude for state-sponsored cyber operations, and has advocated for strong export controls on the most powerful chips used to train AI, all to preserve the U.S. lead in frontier AI development.

25. We have been aligned with the U.S. government’s priority to sustain and enhance America’s global AI dominance to promote economic growth, human flourishing, and national security.

26. Since as early as 2024, Anthropic has led the field in supporting U.S. national security priorities. We collaborated closely with national security stakeholders on a variety of initiatives to advance our shared goal of building safe AI systems. These include collaboration with federal partners on AI safety research, evaluation frameworks, and strategic cloud partnerships as AI assumed a more prominent national security role. As a result, Anthropic’s AI models were the first ever to be used by American warfighters on classified systems. Today, Claude is reportedly the Department of War’s (“DoW” or the Department”) most widely deployed frontier AI model and the only one currently on classified systems.

27. We were also the first to proactively work to align our AI system with the government’s national security needs. Anthropic developed Claude Gov, a dedicated model for national security users to address real-world operational needs that also includes a government-specific addendum to the Usage Policy described above. While Claude Gov underwent the same

1 rigorous safety testing as all other Claude models, it was designed to fulfill the missions of our
2 national security customers and is more likely to comply with requests that are appropriate in a
3 military context. For example, some standard versions of Claude refuse to analyze documents that
4 appear to be classified, such as materials marked Top Secret. That restriction is appropriate for
5 ordinary commercial users, but it would be incompatible with legitimate national security uses by
6 government personnel, and Claude Gov will not refuse to analyze Top Secret documents.

7 28. The government-specific Usage Policy addendum was designed to strike a balance
8 between enabling national security beneficial uses and mitigating potential harms. For example,
9 whereas ordinary civilians do not conduct foreign intelligence analysis, the government's military
10 and intelligence communities do. As a result, the government-specific addendum does not impose
11 the same restrictions on national security use as it does on civilian customers. In our recent
12 negotiations discussed below, for example, we made clear that we were prepared to authorize use
13 of our models for additional purposes for the DoW—such as developing more effective
14 weapons—but we have never allowed regular users to use Claude to assist with weapons
15 development.

16 29. Anthropic partnered closely with national security prime contractors to enable the
17 provision of our models on their platforms, through which DoW and other government national
18 security customers access AI systems. DoW gained access to Claude Gov for the first time in
19 March 2025 via Anthropic partner platforms—and its usage was governed by Anthropic's Usage
20 Policy and the government-specific addendum. During this period, the government-specific
21 addendum imposed broad restrictions that would have prohibited mass surveillance of Americans
22 and lethal autonomous warfare. DoW assented to these terms by using Claude Gov through the
23 platforms of Anthropic's partners and, to my knowledge, did not object to them at any point.

24 30. In July 2025, Anthropic engaged in a separate negotiation with DoW regarding how
25 the Department might further expand its usage of our models by accessing them directly from
26 Anthropic rather than through our partners. These discussions never advanced beyond scoping out
27 potential work; because we did not reach the implementation phase, the terms of a Usage Policy
28 were not discussed.

1 31. Meanwhile, throughout this period, DoW continued to use our models through
2 partner platforms, subject to the broad restrictions of the government-specific addendum to the
3 Usage Policy described above. DoW was satisfied with our models, embedded them into its
4 operations, and expanded their usage.

5 32. In the fall of 2025, DoW and Anthropic began negotiations regarding a new
6 deployment of our models on DoW's GenAI.mil platform. The discussion contemplated various
7 types of deployments, including some that, if implemented, might require a direct contractual
8 relationship between DoW and Anthropic, including with respect to the Usage Policy. During
9 these negotiations, DoW asked Anthropic to remove its Usage Policy not just with respect to the
10 GenAI.mil platform but across all existing and future offerings and to permit DoW, and its
11 contractors and subcontractors, to use all versions of Claude for "all lawful uses." Anthropic
12 engaged in these negotiations in an effort to support DoW's national security priorities in a
13 manner consistent with the company's core principles. As part of these efforts, the Department
14 sent partial contract language incorporating this term to Anthropic and delivered an ultimatum:
15 Anthropic must agree to the revised term or lose all current and future Department business.
16 Contract modifications for facility and personnel clearances and classified work have been frozen
17 since then as the parties continue to discuss the Department's demand.

18 33. Anthropic ultimately agreed to allow DoW, and its contractors and subcontractors,
19 to use Claude without a broad set of restrictions that had previously applied to all DoW usage.
20 However, Anthropic set two critical exceptions: mass surveillance of Americans and lethal
21 autonomous warfare. With those two limitations, Anthropic agreed to "all lawful uses" of Claude.
22 This change, if accepted, would have shifted the structure of the government-specific addendum
23 from a "whitelist" approach—under which broad prohibitions applied with limited authorized
24 exceptions—to a "blacklist" approach, under which all lawful uses were permitted except for these
25 two prohibitions.

26 34. First, we would not agree to the use of Claude to carry out mass surveillance of
27 Americans. It is my understanding that existing surveillance laws were written before the advent
28 of frontier AI systems. Tools like Claude enable aggregation and analysis of massive datasets at

1 unprecedented scale, potentially facilitating practices inconsistent with Americans’ rights even if
2 they appear arguably compliant with laws written before the advent of AI and interpreted by
3 courts only in a pre-AI context. For example, while there is generally no expectation of privacy in
4 public spaces, powerful AI could enable the government to aggregate and analyze millions of
5 public surveillance camera feeds into real-time, population-scale tracking—capabilities not
6 contemplated or addressed by existing federal law. Our legal frameworks have not yet adapted to
7 these novel technologies. Especially in a moment where technology has so outpaced legal
8 frameworks, we at Anthropic, based on our distinctive understanding of what this technology can
9 effectuate, do not believe it is safe or responsible for an AI developer to knowingly enable
10 large-scale surveillance of Americans. Permitting such use would risk Claude being misused in
11 ways that seriously infringe Americans’ rights. Moreover, removing these limitations would also
12 create a risk of inadvertent harm, such as by Claude collecting more information about U.S.
13 persons than the user intended. For example, a user might ask Claude to obtain a specific piece of
14 information about a U.S. person that is lawful to query, but because Claude is operating in an
15 unsafe context, it could inadvertently collect or synthesize a far broader set of information about
16 that individual—including information that the U.S. government is not permitted to collect or
17 query for certain purposes—even absent any intent by the user to violate rules designed to
18 safeguard civil liberties.

19 35. Second, we would not agree to the use of Claude for lethal autonomous warfare.
20 Lethal autonomous warfare consists of using AI to control weapons without any human oversight
21 when human lives are at risk. Such applications include, for example, an AI-controlled aerial
22 system that independently identifies and classifies an object as a military target, determines
23 engagement criteria are satisfied, and launches a weapon strike without any human reviewing,
24 approving, or having the ability to override decisions made by AI. In our view, today’s AI
25 systems—including Claude—are not capable of reliably carrying out lethal autonomous warfare;
26 this is why we have insisted on meaningful human oversight. As anyone who has used a
27 generative AI tool knows, Claude can make errors. In the context of military options, these errors
28 could have grave consequences, jeopardizing the success of military operations or potentially

1 costing the lives of American soldiers or innocent civilians. For example, it is possible that an AI
2 system could misidentify an American soldier as a terrorist operative. Using Claude in this manner
3 would place America's warfighters and innocent civilians at unacceptable risk.

4 36. Because we trained, and have extensively red-teamed our models, we have a unique
5 understanding of Claude's capabilities and therefore have a deep technical understanding of its
6 limitations. From the perspective of that expertise, we have emphatically concluded that Claude is
7 not yet safe for those uses.

8 37. To be clear, we will not and have never second-guessed the government's national
9 security judgments or missions. Anthropic fully respects that any decisions about military
10 operations rest with DoW. Anthropic also fully respects that because of Claude's limitations and
11 safeguards, DoW may opt to work with another company that better suits its needs. Anthropic has
12 not sought, and would not seek, to dictate how the government conducts its missions and who it
13 works with.

14 38. At the same time, acceding to DoW's demand that we remove these two policy-layer
15 safeguards limitations would undercut Anthropic's core identity and competitive advantage.
16 Anthropic has built its identity, reputation, and trust with customers, partners, investors, and the
17 public on a principled commitment to safety. Stripping away those safeguards would erode
18 internal and external trust, weaken the company's culture, and threaten its ability to attract and
19 retain the expertise and commitment necessary to build innovative, cutting-edge AI systems—
20 harms that extend well beyond the immediate technical effects of lifting the two use restrictions at
21 issue here.

22 39. Maintaining these restrictions on Claude's use in military operations is essential to
23 Anthropic's mission to advance the safe and beneficial development and use of AI. That is why
24 we articulated these two critical use limitations to DoW: given the current state of our systems,
25 allowing Claude to be used for mass surveillance of Americans or for lethal autonomous warfare
26 would not only contravene our expert technical judgment but also the very principles on which
27 Anthropic was founded.

28 * * *

1 I declare under penalty of perjury, pursuant to 28 U.S.C. § 1746, that the above is true and
2 correct to the best of my knowledge.

3 Executed on June 10, 2026.

/s/ Jared Kaplan

Jared Kaplan
Co-Founder, Anthropic

ATTESTATION PURSUANT TO CIVIL LOCAL RULE 5-1(i)(3)

Pursuant to Civil Local Rule 5-1(i)(3), I attest that concurrence in the filing of this document has been obtained from the other signatories.

By: /s/ Michael J. Mongan
Michael J. Mongan

MICHAEL J. MONGAN (SBN 250374)
michael.mongan@wilmerhale.com
WILMER CUTLER PICKERING
HALE AND DORR LLP
50 California Street, Suite 3600
San Francisco, CA 94111
Telephone: (628) 235-1000

EMILY BARNET (*pro hac vice*)
emily.barnet@wilmerhale.com
WILMER CUTLER PICKERING
HALE AND DORR LLP
7 World Trade Center
250 Greenwich Street
New York, NY 10007
Telephone: (212) 230-8800

Attorneys for Plaintiff Anthropic PBC

KELLY P. DUNBAR (*pro hac vice*)
kelly.dunbar@wilmerhale.com
JOSHUA A. GELTZER (*pro hac vice*)
joshua.geltzer@wilmerhale.com
KEVIN M. LAMB (*pro hac vice*)
kevin.lamb@wilmerhale.com
SUSAN HENNESSEY (*pro hac vice*)
susan.hennessey@wilmerhale.com
LAUREN MOXLEY BEATTY (SBN 308333)
lauren.beatty@wilmerhale.com
MATTHEW E. MORRIS (*pro hac vice*)
matt.morris@wilmerhale.com
BARDIA VASEGHI (*pro hac vice*)
bardia.vaseghi@wilmerhale.com
DAKOTA C. FOSTER (*pro hac vice*)
dakota.foster@wilmerhale.com
WILMER CUTLER PICKERING
HALE AND DORR LLP
2100 Pennsylvania Avenue NW
Washington, DC 20037
Telephone: (202) 663-6000

**UNITED STATES DISTRICT COURT
NORTHERN DISTRICT OF CALIFORNIA**

ANTHROPIC PBC,

Plaintiff,

v.

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

Defendants.

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

DECLARATION OF PAUL SMITH

1 I, Paul Smith, pursuant to 28 U.S.C. § 1746, declare as follows:

2 1. My name is Paul Smith. I am the Chief Commercial Officer (“CCO”) of Anthropic
3 PBC (“Anthropic”), where I have worked since 2025. As Anthropic’s CCO, I am responsible for
4 driving the trusted enterprise adoption of Anthropic’s AI systems and for maintaining long-term
5 customer confidence to deploy Anthropic’s large language model (“LLM”), Claude, in their
6 highly-regulated and business-critical environments.

7 2. Prior to joining Anthropic, I spent the last thirty years in senior commercial
8 leadership roles at major enterprise technology companies, focusing on building and scaling their
9 global go-to-market strategies. Most recently, I served as the President of Global Customer and
10 Field Operations at ServiceNow, where I oversaw worldwide sales and customer operations. Prior
11 to that role, I held senior leadership roles at Salesforce and Microsoft.

12 3. As the CCO for Anthropic, I have personal knowledge of the contents of this
13 declaration, or have knowledge of the matters based on my review of information and records
14 gathered by Anthropic personnel, and could testify thereto.

15 4. I understand that Anthropic and the Department of War (the “Department”) had
16 been discussing a direct agreement to deploy our AI models on classified systems. Despite
17 ongoing negotiations, on February 27, 2026, President Trump posted on Truth Social, directing all
18 agencies to cease use of Anthropic’s AI system immediately. Secretary of War Pete Hegseth then
19 posted on X.com that he was designating Anthropic “a supply chain risk,” which he claimed
20 would prohibit contractors, suppliers, and partners that conduct business with the United States
21 military from engaging commercially with Anthropic. On March 4, 2026, Anthropic received two
22 letters signed by Secretary Hegseth and dated March 3, 2026, that notified the company of its
23 purported designation. I refer to these actions collectively below as the “Government’s Actions.”

24 **The Government’s Actions And Statements Have Tarnished Anthropic’s Reputation**

25 5. The Government’s Actions attempt to blacklist Anthropic. They send a clear
26 message to the market: do not associate with Anthropic, or else. The Government’s Actions signal
27 that Anthropic is *persona non grata* in the eyes of the government, a message that threatens to
28 reverberate across the broader market. Like any complex, large business, Anthropic depends on an

1 interconnected network of relationships—including with government agencies, prime contractors,
2 commercial partners, cloud providers, investors, current or prospective employees, customers, and
3 the public. Reputation underpins each of these relationships, and harm to Anthropic’s reputation in
4 one context inevitably damages its reputation in others.

5 6. Anthropic has highly prioritized enterprise adoption of Claude, and the majority of
6 our revenue comes through enterprise customer contracts. Among its offerings, Anthropic makes
7 its Claude models available to customers via an application programming interface—also called an
8 API—for customer integration into their own products. Some customers in turn sell these Claude-
9 integrated products to the U.S. government. In fact, our customers sell Claude-integrated products
10 and services that span the economy in industries including healthcare, education, financial
11 services, manufacturing, retail, software, and energy, just to name a few.

12 7. The unfounded designation of Anthropic as a purported “supply chain risk” is a
13 direct attack on the company’s reputation. I understand that the statutes governing “supply chain
14 risk” relate to foreign adversaries that may harm U.S. national security and, specifically, pose a
15 risk to the Department’s information systems, not to U.S. companies that do not present such a
16 threat.

17 8. The practical effect of the Government’s Actions is to brand Anthropic as akin to a
18 foreign adversary and, in turn, to signal to all American companies that, if they work with
19 Anthropic, they have chosen an unacceptable counterparty. This designation, should it stand,
20 carries a powerful stigma and risks Anthropic’s partnerships with all kinds of firms, not only ones
21 that do business with the government.

22 9. The manner in which the government has claimed to designate Anthropic a supply
23 chain risk significantly exacerbates the risk of reputational harm to the company. The President’s
24 social media post exclaimed that Anthropic’s employees were “Leftwing nut jobs” whose
25 “selfishness is putting AMERICAN LIVES at risk, our Troops in danger, and our National
26 Security in JEOPARDY.” The post further labeled Anthropic as an “out-of-control, Radical Left
27 AI company run by people who have no idea what the real World is all about.” Secretary
28 Hegseth’s statements echoed these false and derogatory claims, asserting that Anthropic had

1 “attempted to strong-arm the United States military into submission - a cowardly act of corporate
2 virtue-signaling that places Silicon Valley ideology above American lives,” and characterizing
3 Anthropic’s position on permissible uses of Claude as “fundamentally incompatible with
4 American principles.”

5 10. The accusations at the heart of the Government’s Actions do not reflect in any way
6 what Anthropic does or who we are. Yet these statements reinforce a false narrative that Anthropic
7 is untrustworthy and unpatriotic, and they endeavor to signal to customers and counterparties that
8 we are a company to be avoided.

9 **The Government’s Actions, If Allowed To Stand, Will Have Significant**
10 **Consequences On Anthropic’s Business Partnerships**

11 11. The Government’s Actions will have far-reaching detrimental effects on Anthropic’s
12 partnerships, not only with Department contractors that currently partner with Anthropic, but also
13 with other Anthropic partners that work with different components of the executive branch and,
14 more broadly, with commercial partners whose interests are unrelated to the national security
15 sector.

16 12. First, the designation has already harmed our relationships with the Department’s
17 contractors, which have been integral to our development as a frontier AI lab supporting U.S.
18 national security objectives. Anthropic was founded and has been led with the core mission of
19 developing safe and beneficial AI that can be used to advance U.S. national security. The
20 Government’s Actions put that mission at risk. Over the last several years, our company has
21 invested significant time and resources in building trust with national security stakeholders and
22 demonstrating that our models can effectively serve the military and intelligence community. The
23 government is now demanding that Anthropic sever those relationships entirely, instantly
24 unsettling carefully cultivated partnerships in the national security arena. As just one example,
25 Anthropic has maintained a multi-year relationship with a defense technology provider that
26 permits Anthropic to offer Claude to Department end-users working on classified datasets.
27 Following the Government’s Actions, that defense technology provider has indicated it intends to
28 move all U.S. government work to other generative AI model providers as soon as possible.

1 13. Second, the Government’s Actions are likely to impair Anthropic’s ability to engage
2 with other components of the federal government and their contractors. On the same day President
3 Trump and Secretary Hegseth issued their directives, the General Services Administration
4 (“GSA”) announced that it was removing Anthropic from USAi.gov, the centralized platform for
5 federal agencies to access and adopt AI tools. That decision cuts off Anthropic’s government
6 procurement opportunities with numerous federal entities, including the U.S. Departments of
7 Veterans Affairs, Health and Human Services (“HHS”), State, Labor, Interior, as well as the
8 federal judiciary. A few days later, on March 2, 2026, Secretary of Treasury Scott Bessent
9 announced on X.com that the Department of Treasury would terminate all use of Anthropic’s AI
10 models in compliance with President Trump’s directive. The Department of State and HHS
11 subsequently issued internal statements announcing the same. The Lawrence Livermore National
12 Laboratory, a nuclear weapons research and development center funded by the Department of
13 Energy, likewise informed Anthropic that it was shutting down Claude, expressing hope that it
14 might one day be permitted to return to the facility.

15 14. Worse still, government contractors for other components of the federal government
16 have been directed to cut ties with Anthropic and to stop using Claude. One partner, which has a
17 multi-million-dollar annual contract, immediately switched from Claude to a competing
18 generative AI model for a deployment by their end customer, the U.S. Food and Drug
19 Administration. That switch instantly eliminated an anticipated revenue pipeline worth more than
20 one hundred million dollars. The government told one electronics testing firm that it must stop
21 using Anthropic for any work related to its government contract. A software security company
22 was likewise instructed by its government client to immediately terminate access to Anthropic’s
23 AI models. Despite acknowledging that there was no legal basis for the directive—only political
24 pressure—the company stated that it had no choice but to comply. These developments make clear
25 that Anthropic’s relationships with entities outside the national security sector are already
26 beginning to erode as a result of the government’s pressure. This harm will only intensify with the
27 issuance of Secretary Hegseth’s recent letters that attempt to implement his February 27 directive
28 posted on X.com.

1 15. Third, based on my understanding, the Government's Actions attempt to expand
2 their impacts beyond Anthropic's ability to provide services to the Department—or even other
3 government agencies—by sowing doubt and uncertainty across Anthropic's commercial
4 partnerships more broadly. In doing so, the Government's Actions cast a long shadow of
5 uncertainty across our business, eroding Anthropic's carefully cultivated relationships with
6 companies in industries including healthcare, education, financial services, manufacturing, retail,
7 software, and energy, just to name a few, and causing our customers to question whether they can
8 engage with Anthropic in any commercial context.

9 16. Indeed, this risk is already materializing. Even before the March 4 letters, in
10 response to the government's February 27, 2026, public threats, many large enterprise customers
11 signaled that publicly doing business with Anthropic had become more costly than working with
12 Anthropic's competitors. This reluctance has taken several forms, including delays in contract
13 discussions; customers and prospects declining to continue evaluating Claude alongside
14 competing LLMs; cancelled sales meetings; withdrawal from planned co-marketing efforts and
15 demands for new contractual protections. For example, our negotiations with three leading
16 financial services institutions have been impacted since the Government's Actions, one of which
17 is valued at one hundred million dollars and had been on the verge of closing. Two of the other
18 firms have made clear that they cannot close their anticipated deals, valued together at over eighty
19 million dollars, unless they obtain newly requested provisions allowing for unilateral contract
20 termination. A national grocery chain cancelled a sales meeting, explicitly noting that it needed to
21 assess the business impacts of the Administration's public statements. Customers across industries
22 as varied as healthcare and cybersecurity have also withdrawn from joint press releases.

23 17. Other contract negotiations with multiple companies have likewise become more
24 challenging. One current financial services customer paused negotiations on a contract worth
25 fifteen million dollars while their legal team engaged in supply-chain-risk-designation "diligence."
26 One of the world's largest pharmaceutical companies is seeking to shorten the intended duration
27 of its contract by ten months. A current financial technology customer explicitly tied cutting a \$10
28

1 million contract to \$5 million, noting that “the DoW situation” had made them unwilling to
2 commit to spending more on Claude.

3 18. Since the Government’s Actions, several customers have begun pivoting their
4 evaluation or deployment of LLMs from Claude to models offered by Anthropic’s competitors.
5 These competitive losses include a state higher-education system comprising more than 20
6 schools, a customer in the business-to-business collaboration software industry, and a
7 telecommunications/media customer that is switching its pilot from Claude Code to a competing
8 AI coding tool.

9 19. Consistent with examples provided above, many companies are also now seeking
10 contractual provisions to protect against uncertainty, when they had not previously done so. For
11 example, as previously noted, a household-name financial services company is now stalling
12 negotiations over a deal valued at well over \$50 million, seeking to add a termination-for-
13 convenience provision to its contract. And a legal-technology company in the midst of a million-
14 dollar contract negotiation sought to change the structure of its contract from committed spending
15 to a pay-for-use structure, citing its investment in government relationships and the perceived risk
16 created by the Government’s Actions. Some have described these provisions as a precautionary
17 measure in the event the government orders them to stop working with Anthropic. Other
18 customers requesting these provisions have stated they would intend to exercise their newly
19 obtained termination rights in the event of a formal government directive designating Anthropic a
20 supply chain risk—we expect they will now do so, in light of the Government’s Actions. And,
21 from the beginning of this dispute, all have taken steps that reflect deep distrust and a growing fear
22 of associating with Anthropic while the Government’s Actions loom.

23 20. All told, Anthropic has received inquiries regarding the Government’s Actions from
24 over one hundred enterprise customers expressing deep fear, confusion, and doubt about
25 Anthropic and the repercussions of associating with our company. For example, a multi-billion-
26 dollar software company expressed uncertainty about its ability to continue using Claude because
27 it maintains a shared codebase across work for the Department and other clients. A large customer
28 that spends hundreds of millions of dollars annually with Anthropic asked whether Claude might

1 be removed from its cloud service provider, and indicated it would require significant additional
2 technical work to serve government customers if a broad ban were imposed that prevented its end
3 customers from using Claude. A Fortune 20 company stated that its lawyers were “freaked out”
4 about working with Anthropic because it does significant government business. And the list goes
5 on.

6 21. Importantly, this harm is extremely challenging to reverse if the designation of
7 Anthropic persists. These relationships were built over years through sustained investment, trust,
8 and repeated engagement. If severed, they will be extraordinarily difficult—and in some cases
9 impossible—to rebuild, particularly for companies whose core business depends on government
10 contracting.

11 * * *

12 I declare under penalty of perjury that the above is true and correct to the best of my
13 knowledge.

14 Executed on June 10, 2026.

/s/ Paul Smith

Paul Smith
Chief Commercial Officer, Anthropic

ATTESTATION PURSUANT TO CIVIL LOCAL RULE 5-1(i)(3)

Pursuant to Civil Local Rule 5-1(i)(3), I attest that concurrence in the filing of this document has been obtained from the other signatories.

By: /s/ Michael J. Mongan
Michael J. Mongan

MICHAEL J. MONGAN (SBN 250374)
michael.mongan@wilmerhale.com
WILMER CUTLER PICKERING
HALE AND DORR LLP
50 California Street, Suite 3600
San Francisco, CA 94111
Telephone: (628) 235-1000

EMILY BARNET (*pro hac vice*)
emily.barnet@wilmerhale.com
WILMER CUTLER PICKERING
HALE AND DORR LLP
7 World Trade Center
250 Greenwich Street
New York, NY 10007
Telephone: (212) 230-8800

Attorneys for Plaintiff Anthropic PBC

KELLY P. DUNBAR (*pro hac vice*)
kelly.dunbar@wilmerhale.com
JOSHUA A. GELTZER (*pro hac vice*)
joshua.geltzer@wilmerhale.com
KEVIN M. LAMB (*pro hac vice*)
kevin.lamb@wilmerhale.com
SUSAN HENNESSEY (*pro hac vice*)
susan.hennessey@wilmerhale.com
LAUREN MOXLEY BEATTY (SBN 308333)
lauren.beatty@wilmerhale.com
MATTHEW E. MORRIS (*pro hac vice*)
matt.morris@wilmerhale.com
BARDIA VASEGHI (*pro hac vice*)
bardia.vaseghi@wilmerhale.com
DAKOTA C. FOSTER (*pro hac vice*)
dakota.foster@wilmerhale.com
WILMER CUTLER PICKERING
HALE AND DORR LLP
2100 Pennsylvania Avenue NW
Washington, DC 20037
Telephone: (202) 663-6000

UNITED STATES DISTRICT COURT
NORTHERN DISTRICT OF CALIFORNIA

ANTHROPIC PBC,

Plaintiff,

v.

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

Defendants.

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

DECLARATION OF KRISHNA RAO

1 I, Krishna Rao, pursuant to 28 U.S.C. § 1746, declare as follows:

2 1. My name is Krishna Rao. I am the Chief Financial Officer (“CFO”) at Anthropic
3 PBC (“Anthropic”), where I have worked since May 2024. As CFO, I am responsible for
4 overseeing Anthropic’s financial strategy and operations, including capital allocation, financial
5 risk assessment, investor relationships, capital raising, and ensuring that Anthropic maintains the
6 financial discipline and operational stability necessary to support the responsible development and
7 deployment of safe, advanced artificial intelligence (“AI”) systems.

8 2. Before joining Anthropic, I was the CFO of Fanatics Commerce (a licensed
9 consumer products sports platform) and the first CFO at Cedar (a healthcare payments and patient
10 engagement platform). Prior to those roles, I served as Global Head of Corporate and Business
11 Development and led Corporate and Operations FP&A (Financial Planning and Analysis) at
12 Airbnb. Earlier in my career, I was a private equity investor at Blackstone and a strategy
13 consultant at Bain & Company. I received a J.D. from Yale Law School and an A.B. in Economics
14 from Harvard College.

15 3. As the CFO for Anthropic, I have personal knowledge of the contents of this
16 declaration, or have knowledge of the matters based on my review of information and records
17 gathered by Anthropic personnel, and could testify thereto.

18 **The Uncertainty Created By The Government’s Actions Is Concretely Harming Anthropic**

19 4. Recent actions by the government have created significant uncertainty in the market.
20 For example, over the weekend after the President Trump’s and Secretary Hegseth’s social media
21 posts, a major investor in Anthropic informed me that the Department of War (the “Department”)
22 had contacted several of its portfolio companies about their use of Claude. Those companies have
23 grown worried and uncertain about their ability to use Claude. A different Anthropic investor
24 forwarded me market analysis, which called Anthropic’s supposed designation as a supply chain
25 risk a “sanction” and opined that any company that wants to serve as a federal contractor cannot
26 do business with Anthropic. I have seen multiple client alerts from law firms describing the
27 potentially far-reaching nature of the government’s actions and suggesting that Department
28 contractors may be best served by reevaluating their relationship with Anthropic. Public reporting

1 has been similarly confused. For example, a *Yahoo* article mistakenly asserts that Secretary
2 Hegseth “ordered the Pentagon to bar its contractors *and their partners* from any commercial
3 activity with Anthropic.”¹

4 5. The uncertainty the government’s actions have created is already inflicting a wide
5 range of challenges and problems on Anthropic.

6 6. My team has estimated revenue exposure across a range of potential customer
7 interpretations of the government’s actions. Insofar as customers adopt a narrow reading—where
8 the government’s actions are understood to prohibit only the use of Claude in work performed for
9 the Department—we estimate that hundreds of millions of 2026 revenue is at risk. Insofar as
10 significant swaths of customers’ risk calculations sweep broader—where customers believe that
11 doing business with Anthropic at all would jeopardize their ability to contract with U.S.
12 agencies—the impact is significantly larger and will vary by customer type. Defense contractors
13 and others with financial dependence on the Department are most likely to adopt the maximal
14 interpretation, and we estimate it could reduce revenue from those customers by 50–100 percent.
15 Across Anthropic’s entire business, and adjusting for how likely any given customer is to take a
16 maximal reading, the government’s actions could reduce Anthropic’s 2026 revenue by multiple
17 billions of dollars.

18 7. When it comes to Anthropic’s investors, even if they recognize the true scope of the
19 relevant authorities being invoked by the government, the mere fact of the purported
20 designation—combined with President Trump’s and Secretary Hegseth’s public statements—risks
21 substantially undermining market confidence and Anthropic’s ability to raise the capital critical to
22 train next-generation models and maintain its position in a very competitive race at the AI frontier.
23 Compounding the problem is the threat of mounting fear and uncertainty among customers. These
24 effects reinforce one another: investors withdraw from companies losing customers, and
25 customers avoid companies perceived as struggling to attract capital or sustain growth.

26
27
28 ¹ Jen Judson, et al., *US bars Anthropic products from agencies, contractors for rejecting US military offer*, Yahoo Finance (Feb. 27, 2026), <https://sg.finance.yahoo.com/news/us-bars-anthropic-products-agencies-213459254.html> (emphasis added).

1 I declare under penalty of perjury that the above is true and correct to the best of my
2 knowledge.

3 Executed on June 10, 2026.

/s/ Krishna Rao

Krishna Rao

Chief Financial Officer, Anthropic

ATTESTATION PURSUANT TO CIVIL LOCAL RULE 5-1(i)(3)

Pursuant to Civil Local Rule 5-1(i)(3), I attest that concurrence in the filing of this document has been obtained from the other signatories.

By: /s/ Michael J. Mongan

Michael J. Mongan

MICHAEL J. MONGAN (SBN 250374)
michael.mongan@wilmerhale.com
WILMER CUTLER PICKERING
HALE AND DORR LLP
50 California Street, Suite 3600
San Francisco, CA 94111
Telephone: (628) 235-1000

EMILY BARNET (*pro hac vice*)
emily.barnet@wilmerhale.com
WILMER CUTLER PICKERING
HALE AND DORR LLP
7 World Trade Center
250 Greenwich Street
New York, NY 10007
Telephone: (212) 230-8800

Attorneys for Plaintiff Anthropic PBC

KELLY P. DUNBAR (*pro hac vice*)
kelly.dunbar@wilmerhale.com
JOSHUA A. GELTZER (*pro hac vice*)
joshua.geltzer@wilmerhale.com
KEVIN M. LAMB (*pro hac vice*)
kevin.lamb@wilmerhale.com
SUSAN HENNESSEY (*pro hac vice*)
susan.hennessey@wilmerhale.com
LAUREN MOXLEY BEATTY (SBN 308333)
lauren.beatty@wilmerhale.com
MATTHEW E. MORRIS (*pro hac vice*)
matt.morris@wilmerhale.com
BARDIA VASEGHI (*pro hac vice*)
bardia.vaseghi@wilmerhale.com
DAKOTA C. FOSTER (*pro hac vice*)
dakota.foster@wilmerhale.com
WILMER CUTLER PICKERING
HALE AND DORR LLP
2100 Pennsylvania Avenue NW
Washington, DC 20037
Telephone: (202) 663-6000

**UNITED STATES DISTRICT COURT
NORTHERN DISTRICT OF CALIFORNIA**

ANTHROPIC PBC,

Plaintiff,

v.

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

Defendants.

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

**DECLARATION OF THIYAGU
RAMASAMY**

1 I, Thiyagu Ramasamy, pursuant to 28 U.S.C. § 1746, declare as follows:

2 1. I am the Head of Public Sector at Anthropic PBC. I have held that position since
3 January 2025. Before joining Anthropic, I worked for Amazon Web Services (“AWS”), where I
4 was a Principal Lead for Data, Analytics, and Artificial Intelligence/Machine Learning. In that
5 role, I was responsible for, among other things, implementing Anthropic’s AI models, Claude, for
6 AWS’s public sector customers, including AWS’s deployment of Claude in classified government
7 networks. Together, my prior experience and current role give me a detailed understanding of
8 Anthropic’s relationship with the U.S. government, the technical capabilities of Anthropic’s AI
9 models, how third-party cloud providers deploy Claude to government customers, and how those
10 customers access and integrate Claude in their workflows.

11 2. I have personal knowledge of the contents of this declaration, or have knowledge of
12 the matters based on my review of information and records gathered by Anthropic personnel, and
13 could testify thereto.

14 **Purpose of Declaration**

15 3. I submit this declaration in support of Anthropic’s Motion for Summary Judgment. I
16 have previously submitted multiple declarations in this matter and in related proceedings before
17 the United States Court of Appeals for the District of Columbia Circuit. This declaration
18 consolidates and supplements the information set forth in those prior declarations to provide the
19 Court with a consolidated account of the relevant facts.

20 4. I am aware that the government has asserted that Anthropic poses a supply chain
21 risk to national security, and that government declarants have offered various rationales in support
22 of that designation. While those explanations have evolved over the course of these proceedings,
23 the government’s assertions continue to reflect a fundamental misunderstanding of how
24 Anthropic’s tools are deployed in classified systems and made available to the government, and
25 they contain multiple factual misstatements.

26 5. I submit this declaration to explain 1) how Anthropic builds, trains, and deploys
27 Claude and the extent of the company’s ability to influence the behavior of models before and
28 after deployment; 2) the Department of War’s adoption and use of Claude in classified networks,

1 including through defense contractor platforms; 3) Anthropic's collaborative working relationship
2 with the Department and U.S. national security agencies; and 4) relevant developments following
3 the Department's supply chain risk designation.

4 **Anthropic's Positive Relationship with the U.S. Government**

5 6. Anthropic is a public benefit corporation whose mission is to ensure that the
6 transition to powerful AI benefits humanity. We view partnering with the U.S. government as an
7 important way to achieve that mission. During my time at Anthropic, we have aggressively
8 pursued opportunities to empower the U.S. government to use our signature large language model
9 (LLM), called Claude. Our government customers include the Department of War and agencies in
10 the Intelligence Community.

11 7. Anthropic began working with the U.S. government as early as 2024. In April 2024,
12 Anthropic became the first frontier AI lab to collaborate with the Department of Energy National
13 Laboratories and the National Nuclear Security Administration to evaluate one of Anthropic's
14 models in a Top Secret classified environment to determine how large language models may
15 contribute to or help to address national security risks in the nuclear domain. Anthropic later
16 expanded its partnership with the Department of Energy National Laboratories by deploying
17 Claude to 10,000 scientists at Lawrence Livermore National Laboratory to help bolster research
18 across nuclear deterrence, energy, materials science, and energy security.

19 8. In November 2024, Anthropic expanded its relationship with the U.S. government
20 via a partnership with the software company Palantir Technologies. Anthropic and Palantir
21 partnered to provide intelligence and defense capabilities to U.S. intelligence and defense
22 agencies. That partnership has allowed Claude to be used to support government operations,
23 including rapidly processing large datasets; autonomously completing complex software
24 engineering projects related to offensive and defensive cyber operations and vulnerability
25 detection; supporting military operations; performing intelligence analysis and threat assessments;
26 and handling national security workflows and other mission-critical tasks integral to national
27 security.
28

1 9. We have only deepened our relationship with the Department and the Intelligence
2 Community since then. In June 2025, we announced a custom set of models built for U.S. national
3 security customers, described further below. In July 2025, alongside other frontier AI labs
4 including Google, OpenAI, and xAI, Anthropic was awarded a two-year, up to \$200 million
5 agreement by the Department’s Chief Digital and Artificial Intelligence Office (“CDAO”), the
6 primary office within the Department responsible for integrating and optimizing AI capabilities
7 across the Department. As part of that agreement, we expanded our commitment to work with the
8 Department to explore and prototype frontier AI capabilities that advance U.S. national security.
9 In announcing the award of these contracts to Anthropic and the other frontier AI labs, the CDAO
10 emphasized that it was “leveraging commercially available solutions” to “accelerate the use of
11 advanced AI” in service of the Department’s mission. Anthropic worked diligently under that
12 agreement, scoping out potential ways that the Department could best be served by Claude and
13 related Anthropic professional services. During this period, the Department conveyed to
14 Anthropic that Claude was the best solution for some of the proposals.

15 10. We have also partnered with the General Services Administration (“GSA”), the
16 civilian agency responsible for centralized procurement and shared services across the federal
17 government. In August 2025, Anthropic and GSA announced a first-of-its-kind OneGov
18 agreement to deliver Claude Gov to all three branches of the government—civilian executive,
19 legislative, and judiciary—for a nominal fee of \$1 per agency. As GSA announced at the time,
20 “This trailblazing partnership directly supports the White House’s America’s AI Action Plan and
21 positions the United States as the global leader in government AI adoption, ensuring that the
22 federal workforce can tap into the transformative power of AI to modernize operations, improve
23 decision-making, and deliver better results for taxpayers.”

24 11. Throughout, we have maintained our commitment to supporting the national
25 security of the United States and its allies. I have a team of over 30 individuals who manage
26 relationships with our federal government customers, including national security customers. This
27 team includes individuals with security clearances who drew on their past experience training AI
28 models for the Department and Intelligence Community to help lead the development of

1 Anthropic's Claude Gov models. My team partners closely with dozens of other Anthropic
2 employees across other parts of the company, including Product, Applied AI, Legal, and Policy.
3 And recently, to help the company identify and develop high-impact applications that strengthen
4 U.S. and close allies' capabilities in areas like cybersecurity and intelligence analysis, on August
5 27, 2025, we introduced the Anthropic National Security and Public Sector Advisory Council, a
6 group of leading bipartisan national security and public policy experts.

7 12. As one would expect, our many partnerships with the U.S. government have
8 involved intense security reviews and thorough vetting. Last year, the Department's Defense
9 Counterintelligence and Security Agency granted Anthropic a Top Secret facility security
10 clearance, after an 18-month vetting process, along with several personnel clearances for
11 Anthropic employees and management, to enable continued support for classified national security
12 projects.

13 13. In June 2025, GSA and the Department granted Claude authorization through the
14 Federal Risk and Authorization Management Program ("FedRAMP") for use with FedRAMP
15 High and DoD Impact Level 4 and 5 workloads, representing the highest levels of cloud security
16 certification for unclassified and controlled unclassified information.

17 14. To my knowledge, at no point in any of those processes did the Department identify
18 any potential supply chain risk posed by Anthropic, its employees, or its products and services.

19 15. In fact, to my knowledge, we have only ever received positive feedback about
20 Claude's performance from our governmental customers. For example, a technology leader at a
21 large civilian agency informed us that their agency had used Claude to resolve legacy system
22 issues that had been stuck for years (with one five-year-old bug fixed within days), build internal
23 tools within days rather than waiting through year-long procurement processes, and modernize
24 applications that have not been updated since the 2000s to current technology stacks within hours.
25 The leader of another organization shared that they named Claude their "top model" and planned
26 to grow usage as quickly as possible: "We want to move as fast as possible with you guys."
27 Similarly, senior leaders within a part of the Intelligence Community reported that they were
28 "hammering away" at Claude once they obtained access to their networks.

Background On Claude

16. Like other LLMs, Claude can interpret and respond to diverse user prompts, including analyzing large volumes of text, generating written content, generating and debugging code, and reasoning through complex, multi-step problems. It can also be configured to act “agentially” by taking actions on a user’s behalf—from simple tasks like sending emails to complex workflows executed with limited ongoing user direction.

17. LLMs like Claude are dual-use technology: the same capabilities that drive medical breakthroughs, accelerate scientific research, and enhance human creativity can also be misused by malicious actors, making it difficult to restrict harmful applications without undermining beneficial ones. LLMs can also produce outputs that diverge from their users’ intentions or reflect errors embedded in training data.

18. To address these risks, Anthropic has implemented a multilayered safety framework operating at three levels: the model layer, the safeguards layer, and the policy layer (with the last effectuated partly through contract terms). The sections below explain how each layer functions in commercial deployments and how each is adapted for national security deployments, where many commercial-use safety mechanisms may not be appropriate. Understanding this framework is essential to understanding when and how Anthropic can influence how the model operates.

Model Layer

19. At the model layer, Anthropic embeds safety directly through training. Training an LLM requires processing large amounts of data to develop a mathematical representation of language. The product of that training process is a set of model weights, which are the numerical parameters that shape how the model responds to inputs by determining how strongly different inputs influence a model’s outputs. In other words, the model weights encode Claude’s learned behaviors and capabilities and govern how Claude responds to any given input.

20. One key training technique is what we call “Constitutional AI,” which trains the model to evaluate and revise its outputs against a set of normative principles, such as balancing being helpful against avoiding harm. We publish “Claude’s Constitution” to be transparent about those principles and how Claude has been trained to behave. We also use techniques like

1 reinforcement learning from human feedback, in which human reviewers rank model outputs to
2 train the model toward responses that better align with human judgment. Through these
3 techniques, Anthropic also embeds guardrails into the model weights, which are specific
4 behavioral constraints that Claude is trained to follow regardless of how it is subsequently
5 deployed or configured. For example, Claude is trained to refuse to generate child exploitation
6 material under any circumstances.

7 21. Anthropic has adapted its model-layer training to account for the specific needs of
8 national security users. As described in detail below, Anthropic engaged directly with government
9 stakeholders to understand operational needs and develop a version of Claude specifically tailored
10 for national security use. That model, Claude Gov, retains the same baseline alignment and safety
11 testing as commercial models, but is trained to specifically perform in classified environments
12 where users can be trusted with prompting the model in ways that, ordinarily, the model would
13 refuse to be responsive to.

14 22. Model training is the primary mechanism through which Anthropic can influence
15 the behavior of models used by the Department because it is the only technical safety layer that
16 operates before deployment. Once a model is delivered to a third-party cloud provider and
17 deployed inside a government-secure enclave, Anthropic has no ability to access it, alter it, or shut
18 it down. Any subsequent change to the model's behavior requires preparing and delivering an
19 entirely new version of the model, a process that requires the Department's approval and is
20 described in detail below. Anthropic does not engage in ongoing administration of models
21 operationally deployed by the Department.

22 **Safeguards Layer**

23 23. For commercial deployments, Anthropic applies a safeguards layer on top of the
24 model at runtime. This consists of technical controls, including classifiers and probes that detect
25 harmful activity in real time, targeted interventions to reduce harmful outputs, and monitoring
26 systems to identify misuse at scale. Generally speaking, classifiers engage in monitoring functions
27 or intervening functions. A monitoring classifier looks at prompts to assign and record a risk score
28 but does not impact system behavior. An intervening classifier operates like an automated filter

1 that reads Claude's output before it reaches the user—if the output appears to provide instructions
2 for synthesizing a dangerous substance, for example, the classifier triggers an intervention that
3 blocks or rewrites the unsafe response before it is delivered.

4 24. While the safeguards layer is very important in Anthropic's commercial
5 deployments, it is far less relevant when it comes to the Department. From a technical standpoint,
6 because Anthropic does not have any ongoing access to models deployed in the Department's
7 classified environments, it cannot operate or access any classifiers, probes, or other runtime
8 monitoring systems. The Department's cloud provider does deploy a monitoring classifier that
9 records risk scores for prompts (though not the text of the prompt itself), but it does not impact the
10 behavior of the model. Anthropic does not control the monitoring classifier in such operational
11 deployments and cannot access any risk score information. Furthermore, there are no intervening
12 classifiers in use in Department deployments.

13 **Policy Layer**

14 25. The policy layer consists of Anthropic's Usage Policy, which governs how Claude
15 may be used. Claude is available to any user only under Terms of Service that incorporate the
16 Usage Policy. The policy defines the scope of permissible use and steers users away from
17 applications for which Claude has not been developed or is not ready. By establishing reasonable
18 limitations on the appropriate uses for our products and services, the Usage Policy effectively
19 defines Anthropic's commercial offerings. For commercial users, the Usage Policy generally
20 prohibits uses that pose unacceptable risks, including using Claude to conduct surveillance, to
21 compromise computer systems, or to design weapons or other systems intended to cause harm or
22 loss of human life. The Usage Policy is critical because technical safeguards alone cannot prevent
23 all dangerous uses of such a powerful, novel technology.

24 26. The contractual usage restrictions at the policy layer are both substantively and
25 technically different for Department deployments. Substantively, Anthropic has adapted its Usage
26 Policy over time, in consultation with government partners, specifically to support national
27 security uses. When Claude was first deployed for national security use in early 2024, the
28 applicable Usage Policy was identical to the commercial policy. In June 2024, Anthropic issued its

1 first government-specific addendum to allow analysis of foreign intelligence; a subsequent
2 January 2026 addendum further expanded permissible government uses to include offensive cyber
3 operations and certain intelligence collection. Technologically, Anthropic does not have the ability
4 to enforce the Usage Policy in Department operational deployments or direct visibility into how
5 the model is being used. However, these usage restrictions are nevertheless critical to articulate the
6 boundaries of safe usage.

7 27. In summary, as a practical matter, Anthropic can influence the operation of its
8 models with respect to the Department's uses only through model training prior to operational
9 deployment. And it can negotiate contract-based but technically unenforceable limitations on use.
10 In both contexts, Anthropic has sought to be exceptionally transparent. It has presented its usage
11 limitations directly to the Department and has worked collaboratively with the government in
12 adjusting model-layer safeguards for national security purposes. Below I describe model training
13 for national security purposes and the Department's rigorous testing, evaluation, and approval
14 process to assess the security, capability, and behavior of models prior to operational deployment.

15 **The Government's Adoption of Anthropic's Models in Classified Environments**

16 28. To understand the ways in which Anthropic has used model training to shape the
17 operation of its models in Department deployments, it is useful to step back to explain the
18 evolution of the Department's operational deployments of Anthropic tools.

19 29. Anthropic began working with the federal government as early as 2024. Beginning
20 in early 2024, Anthropic worked with a third-party cloud provider to support deployment of
21 Anthropic's commercial AI models in classified cloud environments. By May 2024, Claude
22 Sonnet was available on a third-party provider's Top Secret cloud and was being used by
23 Intelligence Community ("IC") and Department customers. As government users adopted the
24 commercial version of Claude in classified settings, Anthropic received feedback that its
25 commercial models, appropriately trained to decline discussion of classified materials with the
26 general public, sometimes applied those same protections to IC and Department users in secure
27 environments. In response, in order to better assist the national security components of the U.S.
28 government in using Claude to fulfill their mission, Anthropic engaged directly with government

1 stakeholders, including through classified discussions led by Anthropic engineers holding security
2 clearances, to understand operational needs and identify appropriate technical solutions.

3 30. In December 2024, the updated commercial version of Claude (Sonnet 3.5) was
4 launched in classified cloud environments. To be clear, in these arrangements, once Anthropic
5 gives its models to the third-party cloud provider, Anthropic has no access to, or control over, the
6 model as deployed or used by government customers.

7 31. In early 2025, Anthropic—on its own initiative and at its own expense—assigned
8 cleared personnel to develop Claude Gov, a version of Claude specifically tailored for national
9 security use and designed to operate in a high-trust classified environment. As discussed above,
10 Claude Gov retains the same baseline alignment and safety testing as commercial models but is
11 trained to understand what entity is making use of it and what legal framework governs that work.
12 Specifically, Claude Gov is instructed to assume that it is working with a cleared U.S. military or
13 intelligence officer acting under lawful authority. Claude Gov was developed based on extensive
14 collaboration with government partners regarding applicable policies and operational use cases. It
15 is trained to perform tasks appropriate to the national security context while remaining within the
16 rules governing intelligence and military activities.

17 32. Claude Gov is a distinct version of the Claude model, specifically trained to operate
18 appropriately in a national security context. Claude Gov does not include any additional
19 “wrapper” layer that could provide Anthropic with access to, or control over, the model’s
20 operation or that inserts any additional safeguards or restrictions.

21 33. Claude Gov was made available to IC and Department customers through classified
22 cloud environments in March 2025. Government users reported that it performed significantly
23 better for classified uses than the commercial version, including now completing tasks that were
24 appropriate to the national security mission—such as reviewing classified documents or
25 supporting military planning—which it had previously refused. Adoption of Claude Gov increased
26 rapidly across the Department and the IC. Government customers expressed strong appreciation
27 for Anthropic’s proactive approach to national-security partnerships.
28

1 34. In June 2025 and again in November 2025, Anthropic deployed new, improved
2 versions of the Claude Gov models. Before each deployment, Department customers conducted
3 extensive testing to confirm that the new model performed as expected before the prior version
4 was retired and replaced, as described further below. Anthropic closely collaborated with its
5 government and third-party partners to ensure a seamless transition throughout this process.

6 35. Within weeks of the November 2025 release, nearly all IC and Department
7 customers had adopted the updated version of Claude Gov. Government users praised this model
8 as a major technical advance, citing significant improvements in reasoning, agentic tool use, and
9 coding. These improvements enabled customers to fully leverage air-gapped deployment of
10 Claude Code, Anthropic’s agentic coding tool, which is now widely used across the IC and
11 Department.

12 **The Department Tests and Approves All Models Prior to Deployment**

13 36. Anthropic’s role in deploying models on Department systems is limited to providing
14 the model. The Department and its cloud provider conduct a rigorous, multilayered evaluation and
15 approval process, entirely within their own systems and under their own control.

16 37. Before any model can be deployed in a classified environment, the third-party cloud
17 provider that operates the relevant government region must submit a Significant Change Request
18 (“SCR”) to the relevant government region sponsor. The SCR process is a formal approval
19 mechanism by which the cloud provider certifies that a proposed system change—including the
20 introduction of a new model version—meets applicable security and operational requirements.

21 38. SCRs are classified according to complexity levels. A low-complexity SCR involves
22 straightforward, well-understood changes with minimal risk and typically requires the fewest
23 approvals and shortest review period. A medium-complexity SCR involves changes with a
24 moderate degree of technical risk or interdependency and requires additional review. A high-
25 complexity SCR involves significant system changes with broad potential impact and requires the
26 most comprehensive review and approval process.

27 39. The transition of a new model from the unclassified environment (the “low side”) to
28 the classified environment (the “high side”) is controlled entirely by the cloud provider.

1 Anthropic's involvement ends at the low side. Anthropic provides the model to the cloud provider;
2 thereafter, the cloud provider takes responsibility for validating, transferring, and deploying the
3 model in the government's secure environment. Anthropic personnel are not present during or
4 after that transition and have no access to or visibility into the high-side system.

5 40. Once a new model version is available in the high-side environment, the cloud
6 provider conducts its own security testing before the model may be made available to the
7 Department. That security testing is separate from and in addition to Anthropic's own internal
8 evaluations conducted prior to releasing the model to the cloud provider.

9 41. The Department also conducts its own independent evaluation of each model at the
10 program level before approving it for operational use. That evaluation is conducted by Department
11 personnel and is not performed or supervised by Anthropic. The Department's evaluation includes
12 formal testing to assess model performance for the specific operational requirements of the
13 relevant program. Such evaluation includes refusal testing—assessing whether the model correctly
14 responds to tasks appropriate to the program's operational mission—and human-to-model
15 comparisons, in which analysts evaluate model outputs against outputs produced by human
16 analysts for the same prompts, providing a direct measure of model performance relative to
17 established baselines.

18 42. This evaluation process equips the Department to identify and address any concern it
19 may have about potential limitations in a model's performance. The Department has the ability to
20 directly test any aspect of the model's behavior. For example, if the Department were concerned
21 that the model was trained in a way that caused it to refuse tasks the Department deemed
22 appropriate to its lawful mission, or to override the Department's judgment that an activity is
23 permissible, the Department could directly test the model on those specific tasks. Anthropic
24 actively recommends that the Department conduct evaluations and testing before deploying
25 systems in operational environments—particularly for highly consequential use cases—and makes
26 its cleared engineers available to support that testing as needed.

27 43. If a model failed to perform as intended, the Department could assess whether the
28 issue could be resolved—for instance, by refining the prompts provided to the model to enable it

1 to correctly recognize a request as authorized. If the concern could not be resolved, the
2 Department could decline to approve the model, either for the particular activity in question or for
3 Department-wide use. It could also provide feedback to Anthropic and request that Anthropic
4 address the issue in training future models—which Anthropic has done on multiple occasions.

5 44. The evaluation process has matured over the course of successive model
6 deployments. Through its experience with earlier models, the Department has developed a
7 baseline understanding of how each model performs particular tasks and maintains records of
8 example prompts alongside the corresponding outputs produced by both human analysts and the
9 model. Using this baseline, the Department evaluates outputs from new models, identifies
10 instances where the model falls short of expected performance, and determines whether
11 corrections are required. Where corrections are necessary, the Department implements them
12 through prompt refinement—adjusting the instructions provided to the model—and related
13 controls. Once the Department determines that a model is safe, effective, and performs as
14 expected for the specific operational requirements of an individual program, the model is
15 approved for use and the prior model may be retired.

16 45. Each Department component conducts its own testing and is not required to retire
17 the prior version until it is satisfied that the new version meets or exceeds the prior version's
18 performance for its specific mission requirements. For example, some components continued to
19 use Claude Gov Sonnet 3.5 after Claude Gov Sonnet 3.7 was released because they preferred its
20 functionality for their particular needs. By contrast, Claude Gov Sonnet 4.5 (released in November
21 2025) represented a significant enough advance that nearly every Department component adopted
22 it within weeks of release.

23 46. Models are evaluated by testing their behavior in response to inputs. That is how the
24 Department evaluates Anthropic's models, how the industry evaluates models generally, and how
25 Anthropic evaluates Claude. The numerosity of model weights is irrelevant to evaluation and I am
26 not aware of any form of AI model evaluation by the Department or commercial industry that tests
27 a model by examining individual model weights. Doing so would be analogous to testing software
28

1 by examining every individual character of source code—of which there are also many millions.
2 Instead, models are evaluated for their performance and behavior under a range of conditions.

3 47. The Department undertakes comprehensive testing of model behavior before
4 authorizing deployments. While no testing regime could test every possible input or validate every
5 possible execution path, this is not a limitation specific to Anthropic’s models and is inherent to
6 any AI model the Department might deploy, regardless of provider. The same is true of
7 conventional software more broadly—including mission-critical software.

8 **Anthropic Cannot Access or Alter Deployed Models**

9 48. As explained above, once a model is deployed in DoW environments, Anthropic has
10 no ability to access, alter, or shut down the deployed model. Anthropic maintains no back door or
11 remote kill switch, and Anthropic personnel cannot log into a Department system to modify or
12 disable the models during an operation. Anthropic does not have ongoing connectivity to, or
13 control over, model instances running in the Department’s secure networks. Only the government
14 and its authorized cloud provider have access to the deployed system. If the government requires
15 the assistance of any Anthropic employee with respect to the deployed model, it would need to
16 request and facilitate that individual’s access to the Department’s systems directly.

17 49. A deployed model is static. The model weights and guardrails remain fixed unless
18 and until a newer version is deployed to replace it. A model does not degrade or change on its
19 own, and model weights are not automatically modified or updated based on the Department’s
20 usage. The concern that AI models may “drift” or degrade over time—leaving the Department
21 reliant on Anthropic to ensure continued accuracy and fairness—reflects a fundamental
22 misunderstanding of how AI models operate. Once a model is trained and deployed, its parameters
23 remain fixed. Anthropic cannot push undisclosed or unsanctioned changes to a model after the
24 Department has deployed it.

25 50. The undated “Memorandum for the Record” documenting Under Secretary
26 Michael’s analysis of the purported “Urgent Supply Chain Risk” and the March 17, 2026
27 declaration of Under Secretary of War Emil Michael cite the “relatively opaque nature of large
28 language model technology” as a baseline risk requiring heightened trust in Anthropic. It is true

1 that large language models (“LLMs”) are probabilistic, not deterministic, as Anthropic’s Chief
2 Science Officer has explained, and there is some legitimacy to DoW’s concern about the opacity
3 of these systems generally. But that characteristic applies to all LLMs. Compared to other AI labs,
4 Anthropic is uniquely transparent. We maintain a publicly available set of normative principles—
5 “Claude’s Constitution” —a human-readable document that explains the values and rules our AI is
6 trained to follow. Anthropic is also at the forefront of interpretability research aimed at making the
7 behavior of models like Claude easier to understand. If anything, Anthropic presents a lower
8 baseline risk than other AI labs, not a higher one.

9 51. Because Anthropic has no access to Department systems, Anthropic also cannot
10 exfiltrate Department data or conduct surveillance of Department activities. Anthropic does not
11 have access to the Department’s Claude prompts; because Anthropic lacks any access to this
12 customer data, there is nothing that could be exfiltrated or inspected. Any suggestion that
13 Anthropic could engage in “data exfiltration” of Department information is unfounded.

14 52. The government has suggested that Anthropic could alter its models after
15 deployment, without the Department’s knowledge or consent, by modifying model guardrails or
16 model weights—and that Anthropic is capable of making “ongoing operational alterations” to an
17 existing model that could turn “refusals or guardrails on or off.” That is incorrect. Once a model is
18 running inside a government-secure enclave, Anthropic has no ability to unilaterally alter the
19 guardrails or the model weights. Any such change would require Anthropic to produce an entirely
20 new version of the model, which the Department must approve and take affirmative steps to
21 install. The Department retains complete discretion to decline to install any new model, for any
22 reason or for no reason at all.

23 53. For this reason, the term “update” is a misnomer in this context. Unlike conventional
24 software, a model cannot be updated in place through patches or other changes pushed to a
25 running system. The only way to alter a deployed model’s behavior is to replace it with an entirely
26 new model version. This process is analogous to ordering a cake at a restaurant: once the cake is
27 served, the diner cannot adjust the recipe at the table, and the chef cannot reach into the dining
28 room—secretly or otherwise—to change it either. Even if the cake is missing only a teaspoon of

1 baking soda, it must be made again from scratch with the correct ingredients baked in. Anthropic
2 is the chef. If changes to model behavior are desired, whether small or large, Anthropic can
3 prepare a new version of the model with those changes built in. But that new version is not
4 automatically substituted. Before deployment, the third-party cloud provider and Department
5 security reviewers inspect and approve it to ensure it is safe and appropriate. Only after that
6 approval, and only after the customer confirms the new version performs as expected, can it retire
7 the prior model.

8 54. In some instances when a model produces unwanted refusals, the Department can
9 resolve the issue through prompt engineering—that is, by adjusting how a request is phrased so
10 that the model responds more effectively. This is akin to trying a different phrase in a search
11 engine when the initial phrase fails to yield the desired results. It is something all AI users do
12 when a model does not respond to a prompt as desired, and it does not constitute altering the
13 model itself or turning off refusals or guardrails. When basic prompt adjustments are insufficient
14 to resolve an issue, the only option is to train and deploy an entirely new model—which is not a
15 simple “tweak.”

16 55. Every model version is therefore, in fact, a brand-new deployment subject to the
17 same testing and approval processes described above. A new model version cannot be deployed
18 without the approval of the third-party cloud provider and Department security reviewers, and the
19 prior version is not retired until the customer confirms the new version performs as expected.
20 Department components are free to continue using the prior version indefinitely. Neither the
21 Department nor any Department program is compelled to adopt a new model if it is not fully
22 satisfied with the evaluation and testing process, and the existing model will continue to perform
23 without degradation.

24 56. I understand that the government has alleged that AWS provided approximately
25 twelve cleared Anthropic engineers with administrative access to Department systems, allowing
26 them to affect the functioning of the model by pushing updates or changing its functionality,
27 including the ability to remove the model from the system or shut it down altogether. These
28

1 statements are incorrect. AWS has not provided Anthropic or any of its employees with
2 administrative access to the model or with any access to the Department's operational systems.

3 57. Anthropic does employ certain personnel who hold security clearances. Those
4 clearances exist solely to allow these individuals to participate in classified discussions with the
5 Department and others regarding model functionality, to receive feedback, and to provide training
6 to assist the Department in using the model effectively. These individuals do not have access to
7 operational systems. Cleared Anthropic personnel would have access to the Department's
8 operational systems only if the Department itself specifically requested and facilitated such access,
9 and only for the limited purpose of performing functions at the Department's direction.

10 58. In connection with its work supporting government customers, Anthropic has
11 undergone rigorous security vetting and compliance processes. Anthropic holds a DCSA Top
12 Secret facility security clearance, the result of a comprehensive vetting process that took
13 approximately eighteen months to complete. As noted above, in June 2025, Anthropic received
14 FedRAMP High authorization and was certified at DoD Impact Level 4 and Impact Level 5—the
15 highest classification tiers for commercial cloud services used by the Department of Defense. At
16 no point in any of these review and certification processes did any government entity identify
17 Anthropic as a potential supply chain risk or raise any concern about Anthropic's security posture.

18 **The Use of Claude by Government Contractors**

19 59. The processes and technical limitations described above apply equally when defense
20 contractors use Claude Gov in performing Department contracts. Once Claude Gov is deployed in
21 a classified cloud environment, it can be accessed via an application programming interface
22 ("API"), which allows the model to be used for various purposes by cleared defense contractors.
23 Defense contractors may use Claude Gov directly as a chatbot, use it to process their own
24 dataflows, or integrate it into systems and platforms they develop for the military and intelligence
25 community.

26 60. For example, a major defense contractor might embed Claude Gov in an intelligence
27 fusion system used by the military that ingests satellite images and signals intercepts from a large
28 number of sources and integrates them into a unified operational picture. Claude Gov would serve

1 as the analytical layer within such a system. Rather than manually reviewing reports and cross-
2 referencing sources, a military user would be able to query Claude Gov in plain English to request
3 summaries, ask targeted questions, or identify threats. Claude Gov could synthesize incoming
4 data, generate written assessments, and deliver structured, actionable intelligence within seconds.

5 61. Claude Gov is still hosted by third-party cloud providers in these kinds of
6 deployments, and Anthropic would not have ongoing access to the model. Anthropic could not
7 alter the model, change its behavior, or monitor its use. All models—whether accessed directly by
8 Department users or embedded in contractor systems—are subject to the same rigorous testing and
9 evaluation process described above and must be confirmed to perform as expected before being
10 authorized for use.

11 62. Defense contractors might also use Claude Gov to perform work on Department
12 contracts in classified environments without embedding or integrating Claude Gov into any
13 operational system. For example, a defense contractor might use Claude Gov to write a line of
14 computer code for military software or edit a report being drafted as part of a Department contract.
15 In such cases, the outputs generated by Claude Gov are entirely separate from the model itself and
16 exist as ordinary, static data. Once Claude Gov produces an output—such as a line of code or a
17 document—that output has no connection to the model weights, no link back to Anthropic’s
18 systems, and no mechanism by which the model or Anthropic could later access or modify the
19 data.

20 **The Government’s Other Alleged Concerns Are Misplaced**

21 63. I understand that the government has referenced Anthropic deployment issues at the
22 Center for Disease Control and Prevention (“CDC”) as supporting its designation. As I stated in
23 previous declarations, to the extent this incident relates to CDC work of which I am aware, that
24 characterization reflects a misunderstanding. My understanding is that the issue arose from the
25 CDC’s use of a general commercial model whose standard safeguards—appropriately designed
26 for public users—limited certain interactions related to biomedical research to mitigate public
27 safety risks. Once Anthropic became aware of the issue, we worked promptly with our third-party
28 partners and the government to enable use of the model in a manner appropriate to the CDC’s

1 mission and expertise. As noted above, safeguards suitable for commercial deployments may not
2 be appropriate for government uses, and some government deployments require customization or
3 parameter adjustments. That collaborative, iterative process—balancing safety in commercial
4 settings with valid government uses—is a core tenet of Anthropic’s AI safety approach. That is
5 precisely how we addressed the CDC issue once identified, supporting important CDC research
6 while preventing individual users from engineering dangerous biological weapons.

7 64. The government’s declarations also suggest that Anthropic’s employment of foreign
8 nationals increases “adversarial risk.” Anthropic, like many cutting-edge AI companies, employs a
9 globally diverse team of highly skilled researchers and engineers. Anthropic has undergone
10 extensive U.S. Government security vetting in connection with its classified work and takes
11 seriously its obligation to comply with all security and technical requirements for safeguarding
12 sensitive government information. As stated above, to my knowledge, Anthropic is the only AI
13 company where cleared personnel have led the development of AI models tailored to be deployed
14 in classified environments.

15 **The Department’s Supply Chain Risk Designation**

16 65. The government’s supply chain risk designation arose in the context of a broader
17 dispute over usage terms. Beginning in September 2025, as Anthropic was negotiating its
18 deployment on the Department’s genAI.mil platform, the Department began demanding that
19 Anthropic modify its usage policy as a condition of participation. As I have described above,
20 Anthropic’s usage policies do not give Anthropic operational control over government users—
21 they reflect Anthropic’s own principles regarding appropriate use and cannot be enforced against
22 the Department once the model is deployed on Department systems. And of course, if Anthropic’s
23 Usage Policy restrictions do not meet the Department’s needs, the Department is able to use any
24 other AI system that better meets its requirements. Nonetheless, Anthropic indicated it was
25 prepared to modify its usage policy to accommodate the Department’s concerns, and was willing
26 to remove essentially all usage restrictions for Department purposes. Anthropic retained only two
27 narrow carve-outs: restrictions on the use of Claude to develop lethal autonomous weapons
28 systems and to enable the mass domestic surveillance of United States persons. The Department

1 rejected Anthropic’s proposed modifications and insisted on language permitting “all lawful uses.”
2 I understand that the Department has been or is exerting similar pressure on other leading AI labs
3 to agree to similar demands, as reflected in a memorandum Secretary Hegseth issued on January
4 9, 2026. I further understand that many of these labs have agreed to the Department’s preferred
5 usage terms.

6 66. Prior to the supply chain risk designation, neither the Department nor any other
7 government customer had ever identified Anthropic as a supply chain risk, raised supply chain
8 concerns in any security review or procurement process, or expressed any concern about
9 Anthropic’s security posture. To the contrary, Anthropic’s government customers—including
10 those in the Intelligence Community—had consistently expressed strong appreciation for Claude
11 Gov’s capabilities, praised Anthropic as a trusted national security partner, and enthusiastically
12 adopted new model versions. The supply chain risk designation was issued only after the usage
13 term negotiations reached an impasse, and the record does not reflect any identified supply chain
14 concern that preceded or independently motivated the designation.

15 67. In addition, neither I nor anyone else at Anthropic had seen Under Secretary
16 Michael’s memorandum or declaration until the Department of Justice publicly filed them in this
17 matter. Those materials reflect a fundamental misunderstanding of how Anthropic’s tools are
18 deployed in classified systems and otherwise made available to the government, and they contain
19 multiple factual misstatements. Had the DoW disclosed its rationales before designating Anthropic
20 a supply chain risk, I and other Anthropic subject-matter experts could have corrected these errors.

21 68. On February 27, 2026, the government took a series of formal actions that
22 effectively severed Anthropic’s relationship with the Department. A post attributed to President
23 Trump on Truth Social directed all government agencies to cease use of Anthropic’s products.
24 That same day, Secretary Hegseth issued a formal supply chain risk designation identifying
25 Anthropic as a supply chain risk to national security. The accompanying directive stated that,
26 effective immediately, no contractor, subcontractor, or other commercial entity performing work
27 for or on behalf of the Department may conduct any commercial activity with Anthropic, subject
28 to a six-month period for winding down existing contractual obligations.

**The Department's Designation Has Affected Anthropic's
Contracts, Reputation, And Mission**

69. The Department's Chief Digital and Artificial Intelligence Office ("CDAO") awarded Anthropic in July 2025, alongside Google, OpenAI, and xAI, a two-year, up to \$200 million agreement. The award expressly emphasized the Department's interest in leveraging commercially available solutions, and Anthropic worked diligently over the following months to assess and identify ways its services could best support the Department's mission. The Department communicated that it was interested in moving forward with Anthropic's recommendations, and conveyed that Claude was the best available option for certain of its proposed use cases. The arrangement was cancelled after only a few months as a result of the usage term negotiations. Absent those negotiations and the resulting supply chain risk designation, it is almost certain that Anthropic would have continued to develop its relationship with CDAO and received a substantial portion of the full award.

70. As a result of the usage term negotiations and supply chain designation, Anthropic also lost the opportunity to be deployed on genAI.mil, the Department's secure, official generative AI platform to provide its personnel with access to cutting-edge AI models. Anthropic began negotiations with the Department for this deployment in September 2025. As a result of the usage term negotiations and subsequent supply chain risk designation, Anthropic was no longer in consideration to be deployed on the Department's genAI.mil platform. Two other AI vendors have instead deployed their AI models on the Department's platform—one in December 2025 and another in February 2026.

71. In August 2025, Anthropic entered into a first-of-its-kind, government-wide contract with the General Services Administration under the GSA's OneGov program. That contract made Claude available to agencies across all three branches of the federal government for a nominal fee of \$1 per agency, as part of the White House's America's AI Action Plan initiative to expand access to commercial AI capabilities across the government. Anthropic expected the OneGov contract to generate approximately \$100 million in annual recurring revenue in 2026. Following the February 27 government actions, Anthropic was removed from the government's USAi.gov

1 portal and from the Multiple Award Schedule through which the OneGov contract was offered.
2 This decision cut off sales opportunities for many federal agencies, including the U.S.
3 Departments of Veterans Affairs, Health and Human Services (“HHS”), State, Labor, and Interior,
4 in addition to the federal judiciary and many state and local governments that procure through the
5 Multiple Award Schedules. The expected revenue from the OneGov agreement is now entirely at
6 risk.

7 72. Other agencies also took action in response to the February 27 government actions.
8 On Monday, March 2, 2026, the U.S. Department of the Treasury and the Federal Housing
9 Finance Agency (which oversees Fannie Mae and Freddie Mac) announced they were terminating
10 all use of Claude. A technology leader at a federal civilian agency has also informed me that the
11 DoW advised his agency, and is advising all other civilian agencies, to stop using Claude. We also
12 understand that other agencies, including the Department of State and HHS, have issued internal
13 statements saying they will follow the President’s directive. The AI leadership of one portion of
14 the Intelligence Community informed us that they were preparing for “complete detachment” from
15 Claude based on the “directive” they had received. The Lawrence Livermore National Laboratory,
16 a nuclear weapons research and development center funded by the Department of Energy, also
17 informed Anthropic that it was shutting down Claude. Many of these customers—and others,
18 particularly in the Intelligence Community—have continued to express how much they value their
19 partnership with Anthropic and how harmful losing access to our models will be (setting back
20 their work months or even years). They have also expressed, however, feeling like they have little
21 choice in the matter.

22 73. As of March 24, 2026, the following agencies, or sub-agencies thereof, had contracts
23 with Anthropic or used Anthropic’s technology through a third-party provider, but indicated they
24 would be terminating the contracts or terminating uses of Anthropic’s technology after February
25 27, 2026:

- 26 a) U.S. Department of War
- 27 b) U.S. Department of State
- 28 c) U.S. Department of Health & Human Services

d) U.S. Office of Personnel Management¹

e) U.S. Nuclear Regulatory Commission²

f) U.S. Department of Treasury

74. As of March 2026, the following agencies, or sub-agencies thereof, had contracts to use Anthropic's technology or used Anthropic's technology through a third-party provider:

a) U.S. Department of Commerce³

b) U.S. Social Security Administration

c) U.S. Department of Homeland Security

d) Securities and Exchange Commission

e) National Aeronautics and Space Administration

f) U.S. Department of Energy

g) Federal Reserve Board of Governors

h) National Endowment for the Arts

i) Federal Housing Finance Agency

j) U.S. Department of Veterans Affairs

75. The supply chain risk designation has also frustrated Anthropic's contracts with prime contractors that perform services for the Department. For example, Anthropic has maintained a multi-year relationship with a defense technology provider that permits Anthropic to offer Claude to the Department for classified purposes. Anthropic also works closely with a cloud service provider that offers Claude to government agencies, including the Department. Both partners have indicated that, pursuant to Department leadership directives, they need to discontinue using Anthropic as a subcontractor on Department contracts and have begun doing so.

¹ I understand that the Office of Personnel Management sent a message requesting its personnel not to use Claude following the President's social media post.

² Anthropic was informed that the Nuclear Regulatory Commission was no longer deploying Claude, citing the President's social media post.

³ Within the U.S. Department of Commerce, the International Trade Administration has indicated that it was terminating use of Anthropic's technology, whereas the U.S. Patent and Trademark Office and National Telecommunications and Information Administration had contracts with Anthropic or used Anthropic's technology as of March 2026 but have not announced termination actions.

1 Anthropic has invested significant time and resources in building relationships with these national
2 security stakeholders and the supply chain risk designation demands that those relationships be
3 severed.

4 76. The supply chain risk designation has caused both significant financial harm and
5 concrete, ongoing damage to Anthropic's reputation and customer relationships that will be
6 difficult to undo. In terms of financial impact, Anthropic had experienced approximately a
7 fourfold increase in annual recurring revenue ("ARR") from December 2025 to January 2026, and
8 had projected several hundred million dollars in public sector ARR for 2026—a projection that
9 had been revised upward before the designation. The Department's designation and related
10 government actions place an estimated \$150 million or more in immediate Department ARR at
11 risk, including approximately \$50 million attributable to the CDAO contract alone, and expose
12 Anthropic to over \$500 million in total 2026 ARR losses when accounting for the GSA OneGov
13 contract, state and local government customers, and other government-adjacent relationships that
14 have been disrupted. Beyond the direct financial consequences, Anthropic depends on an
15 interconnected network of relationships—with government agencies, prime contractors,
16 commercial partners, cloud providers, and investors—and its reputation underpins each of them.
17 The Department of War's public designation of Anthropic as a supply chain risk signals to
18 enterprise customers and counterparties that Anthropic is a company to be avoided, a harm that
19 becomes increasingly difficult to unwind the longer the designation remains in effect.

20 77. Anthropic is a company that has invested deeply in its relationship with the federal
21 government, and with the IC and the Department in particular. Supporting the national security of
22 the United States and allied democracies has been a deliberate choice—one we have made because
23 it is central to our mission of developing AI for the benefit of humanity. While we are proud of
24 our commitment to our founding values under these difficult circumstances, and are gratified by
25 indications of some public support, we certainly do not welcome this dispute, which we long
26 sought to avoid, including through good-faith negotiations with the Department. While we respect
27 that decisions about the use of AI in military applications belong to the Department, we are
28 harmed by the lost opportunity to continue to support the national security mission directly, and

1 we believe the country and its national security are harmed as well. Notably, we developed these
2 tools specifically for use by the national security components of the federal government at our
3 own initiative and at significant expense. Being branded a threat to the national security of our
4 own country therefore inflicts more than financial harm.

5 **Anthropic's Ongoing Government Engagement**

6 78. Despite the supply chain risk designation, Anthropic has continued to work
7 collaboratively with the U.S. government to promote national security interests, including by
8 supporting the Department's use of Claude during the six-month transition period.

9 79. On March 26, 2026, the Northern District of California issued a preliminary
10 injunction preventing the government from implementing elements of its supply chain risk
11 designation, but did not require the government to continue using any Anthropic products.
12 Following the injunction, although the designation continues to create uncertainty and hesitancy to
13 build long-term partnerships, we have worked with government customers seeking to resume use
14 of Claude to support their restore access as requested.

15 80. We have also worked closely with the U.S. government in connection with the
16 release of Anthropic's new Mythos model. On April 7, 2026, Anthropic publicly announced
17 Project Glasswing, an initiative to deploy Mythos for defensive cybersecurity work in partnership
18 with over fifty major technology companies and critical infrastructure operators. The initiative is
19 designed to apply Mythos's capabilities to defensive purposes before similar capabilities can be
20 adopted by adversaries. Before making Mythos externally available, Anthropic briefed senior
21 officials across multiple government agencies on the model's capabilities, including its ability to
22 identify and help remediate critical vulnerabilities in widely used software systems. Through
23 Project Glasswing, Anthropic extended access to Mythos to select government partners and
24 committed substantial resources, at no cost to national security customers, to efforts to secure
25 critical software infrastructure.

26 81. On April 17, Anthropic CEO Dario Amodei met with senior White House officials,
27 including the Chief of Staff and the Treasury Secretary, to discuss the responsible use of Mythos. I
28 understand that the White House described those discussions as "productive and constructive" and

1 stated that the parties “discussed opportunities for collaboration, as well as shared approaches and
2 protocols to address the challenges associated with scaling this technology.” Since then, Anthropic
3 has remained in frequent communication with representatives of multiple agencies and
4 departments regarding how best to use Mythos to advance U.S. objectives. My team has, over the
5 past several months, worked diligently to support government national security requests and to
6 make clear that we stand ready to assist the government in advancing U.S. national security, as we
7 have always done.

8 82. These engagements reflect Anthropic’s continued commitment to partnering with
9 the U.S. government to strengthen the nation’s cyber defenses and support national security
10 objectives—even as the supply chain risk designation remains in place.

11 * * *

12 I declare under penalty of perjury that the above is true and correct to the best of my
13 knowledge.

14 Executed on June 10, 2026.

/s/ Thiyagu Ramasamy

Thiyagu Ramasamy
Head of Public Sector, Anthropic

ATTESTATION PURSUANT TO CIVIL LOCAL RULE 5-1(i)(3)

Pursuant to Civil Local Rule 5-1(i)(3), I attest that concurrence in the filing of this document has been obtained from the other signatories.

By: /s/ Michael J. Mongan
Michael J. Mongan

MICHAEL J. MONGAN (SBN 250374)
michael.mongan@wilmerhale.com
WILMER CUTLER PICKERING
HALE AND DORR LLP
50 California Street, Suite 3600
San Francisco, CA 94111
Telephone: (628) 235-1000

EMILY BARNET (*pro hac vice*)
emily.barnet@wilmerhale.com
WILMER CUTLER PICKERING
HALE AND DORR LLP
7 World Trade Center
250 Greenwich Street
New York, NY 10007
Telephone: (212) 230-8800

Attorneys for Plaintiff Anthropic PBC

KELLY P. DUNBAR (*pro hac vice*)
kelly.dunbar@wilmerhale.com
JOSHUA A. GELTZER (*pro hac vice*)
joshua.geltzer@wilmerhale.com
KEVIN M. LAMB (*pro hac vice*)
kevin.lamb@wilmerhale.com
SUSAN HENNESSEY (*pro hac vice*)
susan.hennessey@wilmerhale.com
LAUREN MOXLEY BEATTY (SBN 308333)
lauren.beatty@wilmerhale.com
MATTHEW E. MORRIS (*pro hac vice*)
matt.morris@wilmerhale.com
BARDIA VASEGHI (*pro hac vice*)
bardia.vaseghi@wilmerhale.com
DAKOTA C. FOSTER (*pro hac vice*)
dakota.foster@wilmerhale.com
WILMER CUTLER PICKERING
HALE AND DORR LLP
2100 Pennsylvania Avenue NW
Washington, DC 20037
Telephone: (202) 663-6000

UNITED STATES DISTRICT COURT
NORTHERN DISTRICT OF CALIFORNIA

ANTHROPIC PBC,

Plaintiff,

v.

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

Defendants.

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

DECLARATION OF SARAH HECK

1 I, Sarah Heck, pursuant to 28 U.S.C. § 1746, declare as follows:

2 1. I am the Head of Policy at Anthropic PBC (“Anthropic”), where I have worked
3 since June 2024. Before joining Anthropic, I held senior roles at Stripe, the White House’s
4 National Security Council, and the U.S. Department of State.

5 2. In my role as Head of Policy, I lead Anthropic’s public policy engagement,
6 government relationships, strategic partnerships, and government communications, including
7 engagements and initiatives that address the intersection of artificial intelligence (“AI”), national
8 security, and economic policy. This includes being involved in communications between
9 Anthropic and the government.

10 3. As the Head of Policy, I have personal knowledge of the contents of this declaration,
11 or have knowledge of the matters based on my review of information and records gathered by
12 Anthropic personnel, and could testify thereto.

13 4. I have previously submitted two declarations in this matter: a declaration dated
14 March 9, 2026, filed in support of Anthropic’s Motion for a Preliminary Injunction, and a
15 supplemental declaration filed in support of Anthropic’s Reply Brief in Support of Plaintiff’s
16 Motion for a Preliminary Injunction. I submit this consolidated declaration, which sets forth in a
17 single, comprehensive account all of the facts from those two prior declarations, in order to assist
18 the Court.

19 **Anthropic Has Maintained A Strong Relationship With The Federal Government**

20 5. Partnerships with private and public sector entities are central to Anthropic’s
21 mission, business model, public policy goals, and overall success. Our partners range from
22 Fortune 500 companies and U.S. government agencies to small businesses and local
23 municipalities. These partnerships are core to Anthropic and essential to advancing its policy
24 objectives of promoting the responsible deployment of AI, supporting democratic institutions, and
25 ensuring that powerful AI systems are developed and used safely and in ways that benefit society.

26 6. One of Anthropic’s most important partnerships is with the federal government.
27 Anthropic has consistently supported the U.S. government’s goal of maintaining global AI
28 dominance and its efforts to ensure that American AI systems are widely adopted across the

1 federal government and throughout the private sector at home and abroad. Anthropic publicly
2 supported the Trump administration’s efforts to promote AI adoption throughout the federal
3 government as part of its AI Action Plan for America. Anthropic has also partnered with the
4 National Institute of Standards and Technology’s Center for AI Standards and Innovation (CAISI)
5 to undertake collaborative research evaluating and mitigating safety risks. It has advocated in
6 support of the bipartisan Future of AI Innovation Act, which supports CAISI’s initiatives and
7 promotes strong partnerships between private and public stakeholders to advance AI research and
8 innovation.

9 7. Importantly, in developing its relationship with the government, Anthropic has
10 strived to build a reputation as a nonpartisan advocate dedicated to building a safer AI ecosystem.
11 In addition to supporting the bipartisan Future of AI Innovation Act, Anthropic recently donated
12 over \$20 million to Public First Action, a bipartisan nonprofit co-founded and co-led by
13 Republican and Democratic former lawmakers, which supports public education about AI,
14 promotes safeguards, and works to ensure America leads in the AI race. Anthropic has also
15 actively engaged with bipartisan legislative efforts on Capitol Hill, including advocacy in support
16 of the CREATE AI Act of 2025 and the GAIN Act of 2025—both bipartisan safety bills that align
17 with the company’s policy priorities. And the company deliberately maintains a bipartisan
18 lobbying effort, with an in-house team that includes former senior staffers from both sides of the
19 aisle and a roster of Republican-aligned and Democratic-aligned outside lobbying firms.
20 Anthropic invests in and undertakes this public policy work because it is committed to AI safety
21 and the policies that underpin it. Based on my experience, collaboration across the political
22 spectrum is critical to policy progress.

23 8. Anthropic has also been committed to developing its AI systems to advance U.S.
24 national security interests and has partnered with national security components of the government,
25 including the Department of War (the “Department”) and intelligence agencies, to do so. To that
26 end, Anthropic also formed a National Security and Public Sector Advisory Council composed of
27 former U.S. government senior defense and intelligence officials, in part to strengthen the
28 company’s understanding of government needs and become a more effective partner.

1 9. Throughout my engagements with the federal government in my capacity as
2 Anthropic's Head of Policy, various officials and government staff have repeatedly conveyed that
3 the government values its partnership with Anthropic and acknowledged that competitors' AI
4 models lag behind Anthropic's capabilities. A wide range of senior intelligence community
5 officials have emphasized to us that Claude's capabilities are unique, and that Claude is "widely
6 recognized as the leading model for coding and autonomous tool use." Department of War staff
7 themselves have described Claude as "far and away the best model" in briefings attended by the
8 Policy team and warned that losing this capability would set the Department back several years.
9 These characterizations are consistent with the ones set forth in the declarations of Thiya
10 Ramasamy.

11 10. Anthropic appreciates this recognition and has, in turn, repeatedly and publicly
12 expressed its willingness to continue providing Claude for the full range of lawful national
13 security applications, subject to two narrow usage restrictions that reflect the company's expert,
14 considered judgment on the nature of the AI model and its safe and reliable use: namely,
15 restrictions on the use of Claude for mass surveillance of Americans and for lethal autonomous
16 warfare.

17 **Secretary Of War Pete Hegseth Issued An Ultimatum Threatening Consequences If**
18 **Anthropic Adhered To Its Fundamental Commitments**

19 11. Over the months that preceded the initiation of this litigation, Anthropic and the
20 Department had been discussing an agreement to continue their partnership. I understand that the
21 crux of the discussions focused on two specific limitations that would restrict the Department's
22 ability to use Claude for purposes of mass surveillance of Americans and lethal autonomous
23 warfare. During that time, I engaged with the Department on behalf of Anthropic and was
24 included in discussions between Anthropic CEO Dario Amodei and Emil Michael, Under
25 Secretary of War for Research and Engineering and Chief Technology Officer. The conversations
26 throughout this process remained respectful, and the parties remained committed to finding a path
27 forward to harness Anthropic's technology to fulfill the Department's important national security
28 goals.

1 12. Both parties agreed that neither wanted to enter into a begrudging partnership. At
2 one point, Dr. Amodei stressed to Under Secretary Michael that if Anthropic's position on these
3 two guardrails meant that Anthropic was not the right vendor for the Department's needs, then he
4 would respect that decision. Dr. Amodei further emphasized that, if Anthropic and the Department
5 failed to come to an agreement, Anthropic stood ready to assist in an orderly offboarding from the
6 Department's systems.

7 13. On February 24, 2026, I attended a meeting held between Secretary of War Pete
8 Hegseth and Dr. Amodei, among others. At that meeting, Secretary Hegseth said that Claude had
9 "exquisite capabilities." He then objected to Anthropic's unwillingness to remove the two usage
10 restrictions noted above.

11 14. During the meeting, Dr. Amodei reiterated that Anthropic maintains a strong
12 partnership with national security agencies, including the Department and intelligence community.
13 He stated that Anthropic has had no intention of dictating national security operations. Dr. Amodei
14 emphasized that the company's commitment to the safe use of its AI models, as reflected in the
15 agreement, has permitted—and would continue to permit—the Department to use Anthropic's
16 products for all lawful uses, save for those two important exceptions. Dr. Amodei noted that these
17 two exceptions have never obstructed the Department's operations to his knowledge, and
18 Combatant Commanders he has spoken to have been pleased with Anthropic's partnership and
19 models.

20 15. As to autonomous uses, Dr. Amodei clarified that, while Anthropic has been open to
21 certain autonomous applications of its technologies, at this stage, AI systems cannot yet capably or
22 reliably perform lethal autonomous warfare without appropriate oversight.

23 16. Secretary Hegseth stated that Dr. Amodei's concerns were "understandable" and
24 further stated that "they don't do mass domestic surveillance." While Secretary Hegseth expressed
25 that the Department "would love to work with [Amodei]," he stated that the Department had many
26 other vendors to choose from that have never raised these types of concerns and would never hold
27 any veto power over the Department.

1 17. At the end of the meeting, Secretary Hegseth presented Anthropic with what
2 appeared to be an ultimatum: if Anthropic did not agree to “all lawful uses” of its system by 5
3 p.m. on Friday, February 27, 2026, the Department would issue a statement at 5:01 p.m. that they
4 would designate Anthropic a supply-chain risk—preventing Anthropic from partnering with the
5 Department, anyone affiliated with the Department, or any other agency—or invoke the Defense
6 Production Act to compel Anthropic to comply with the Secretary’s demands.

7 18. During this meeting, Secretary Hegseth never stated that Anthropic’s AI models
8 were unsafe, much less subject to compromise by a foreign adversary. I am unaware of anyone at
9 the Department ever suggesting Anthropic’s models are somehow insecure or have been
10 compromised. Moreover, I understand that these two restrictions have never previously been a
11 barrier to the Department’s adoption and use of our models to date. Nor has anyone else
12 articulated such a barrier to me.

13 19. Additionally, it is my understanding that the Department had previously expressed
14 concerns that Anthropic had objected—after the fact—about how Anthropic’s models were being
15 used in Department operations, concerns now reprised in the government’s materials. But this
16 concern was and is misplaced, as Anthropic has made very clear. Indeed, when Secretary Hegseth
17 raised this topic at the February 24 in-person meeting, Dr. Amodei clarified that the company had
18 not objected to the use of its tools as reported. Dr. Amodei further explained that Anthropic has no
19 interest in interfering with the Department’s operations more generally. Dr. Amodei publicly
20 reiterated that stance in a February 26 blog post, explaining that Anthropic “understands that the
21 Department of War, not private companies, makes military decisions” and underscoring that
22 Anthropic has “never raised objections to particular military operations nor attempted to limit use
23 of our technology in an ad hoc manner.”

24 20. Conversations continued even after Dr. Amodei’s meeting with Secretary Hegseth.
25 They remained cordial and amicable. Dr. Amodei continued to seek a resolution prior to Secretary
26 Hegseth’s deadline during the afternoon of Friday, February 27, 2026, during which he provided
27 Under Secretary Michael his redline edits on the Department’s latest offer, alongside a detailed
28 explanation of the same, still attempting to engage in earnest while ensuring that Anthropic’s two

1 guardrails core to its principles would be retained. Under Secretary Michael did not respond in
 2 kind with written edits. At that same time, President Trump directed all agencies, by posting on
 3 Truth Social, to cease use of Anthropic’s AI system immediately. Shortly thereafter, Secretary
 4 Hegseth posted on X.com designating Anthropic a supply chain risk. Throughout these
 5 discussions, Dr. Amodei remained firm in expressing Anthropic’s unmovable guardrails as
 6 essential to the Company’s mission to offer safe, reliable AI tools.

7 **Secretary Hegseth Notified Anthropic Of His “Supply Chain Risk” Designation**

8 21. The following week, as agencies across the federal government moved to implement
 9 the Presidential Directive, Dr. Amodei continued negotiating in good faith with senior Department
 10 officials. Those discussions were still ongoing when, at 8:48 p.m. ET on March 4, Dr. Amodei
 11 received a letter from Secretary Hegseth, expanding on the Secretarial Order’s “supply chain risk
 12 designation.” That letter (the “Secretarial Letter”), dated March 3, asserted that the Department of
 13 War had “determined” that Anthropic’s technology “presents a supply chain risk” and that
 14 exercising the authority granted by 10 U.S.C. § 3252 against Anthropic is “necessary to protect
 15 national security.” The letter pronounced that this determination covers all Anthropic “products”
 16 and “services,” including any that “become available for procurement.” And it asserted that “less
 17 intrusive measures are not reasonably available” to mitigate the risks that Anthropic’s products
 18 and services supposedly pose to national security.

19 22. In the days that followed, the parties continued to work towards an agreement even
 20 while the Department was, unbeknownst to Anthropic, drafting materials claiming the company
 21 poses a risk of sabotage, data exfiltration, and other serious harms. At no point in these
 22 conversations did Under Secretary Michael or anyone at the Department share with Anthropic its
 23 apparently grave concerns or suggest that they were an impediment to a deal. In fact, the
 24 Department and Anthropic were still engaged in negotiations on March 4, 2026—the day after the
 25 Department completed its formalization of the supply-chain risk designation (but before Anthropic
 26 received that formal notice).¹

27
 28 ¹ Attached hereto as **Exhibit 1** is a true and correct copy of a March 4, 2026 email from Emil Michael to Dario Amodei.

The Government's Filings Mischaracterize Anthropic's Positions

23. In reviewing the materials the government has filed in this case, I was struck by its repeated claim that Anthropic insisted during negotiations that the company have some sort of approval role in the Department's operational decision chain. That is false. At no time during Anthropic's negotiations with the Department did I or any other Anthropic employee state that the company wanted that kind of role. We do not want it now, either.

24. Likewise, at no point during the negotiations did any party raise any concerns about Anthropic taking technical steps to disrupt the U.S. military. This is a new purported concern that was raised, for the first time, in the government's filings in this case. Had this concern been raised in a manner allowing Anthropic to respond, we would have explained, among other things, the technical impossibility of such disruption. As I understand from Thiyagu Ramasamy's declaration, Anthropic is unable to alter an AI model after deployment without the Department's knowledge or consent.

25. All of this was avoidable. In an effort to be responsive to the Department's apparent concerns about operational interference, Anthropic was willing to make contractual commitments. In a draft agreement transmitted to the Department on March 4, 2026, Anthropic included language stating: "For the avoidance of doubt, [Anthropic] understands that this license does not grant or confer any right to control or veto lawful Department of War operational decision-making." Anthropic included this language to clarify that it did not seek to be, and would not be, part of any operational decision chain.²

26. I was also surprised to read in the government's materials that Anthropic's beliefs that Claude should not be used for autonomous lethal warfare or mass surveillance of Americans are part of what makes the company a danger to national security. I find this rationale puzzling for two reasons: (1) when negotiations broke down, the parties were very near agreement on language that would address Anthropic's concerns regarding autonomous lethal warfare and (2) Secretary Hegseth himself referred to Anthropic's concerns regarding autonomous lethal warfare and mass

² Attached hereto as **Exhibit 2** is a true and correct copy of a March 4, 2026 email from Anthony Fuscellaro to Dario Amodei Re: Official Correspondence from the Department of War to Anthropic.

1 surveillance of Americans as “understandable.” That was echoed mere days later by the
 2 Commander of U.S. Central Command, who told a media outlet, “Humans will always make final
 3 decisions on what to shoot and what not to shoot.” And a recent threat assessment from the Office
 4 of the Director of National Intelligence, reflecting the collective views of the U.S. Intelligence
 5 Community, further confirms that at least some portions of the Administration believe “it is
 6 essential to make sure that humans maintain control of the [AI models] and how AI is used.”³
 7 Indeed, on March 4, 2026—after the Secretarial Letters were drafted—Under Secretary Michael
 8 emailed Dr. Amodei to say “I think we are very close here” regarding these two issues.

9 27. In recent weeks, Anthropic has worked closely with the U.S. government in
 10 connection with the release of its new Mythos model. On April 7, 2026, Anthropic publicly
 11 announced Project Glasswing, an initiative to deploy Mythos for defensive cybersecurity work in
 12 partnership with over fifty major technology companies and critical infrastructure operators. The
 13 initiative is an effort to put Mythos’s capabilities to work for defensive purposes before similar
 14 capabilities can be adopted by adversaries. Prior to making Mythos available externally, Anthropic
 15 briefed senior officials across multiple government agencies on the model’s capabilities, including
 16 its ability to identify and help remediate critical software vulnerabilities in widely used systems.
 17 Through Project Glasswing, Anthropic extended access to Mythos to select government partners
 18 and committed substantial resources to support efforts to secure critical software infrastructure.
 19 On April 17, Dr. Amodei met with senior White House officials, including the Chief of Staff and
 20 the Treasury Secretary, to discuss the responsible use of Mythos. I understand that the White
 21 House described those discussions as “productive and constructive” and stated that the parties
 22 “discussed opportunities for collaboration, as well as shared approaches and protocols to address
 23 the challenges associated with scaling this technology.”⁴ Since then, we have been in frequent
 24 communication with representatives of various agencies and departments about the best way to

25
 26 ³ Off. of the Dir. of Nat’l Intelligence, *Annual Threat Assessment of the U.S. Intelligence*
 27 *Community* 12 (Mar. 2026), [https://www.dni.gov/files/ODNI/documents/assessments/ATA-2026-](https://www.dni.gov/files/ODNI/documents/assessments/ATA-2026-Unclassified-Report.pdf)
 28 [Unclassified-Report.pdf](https://www.dni.gov/files/ODNI/documents/assessments/ATA-2026-Unclassified-Report.pdf).

⁴ Brendan Bordelon et al., *White House Meets with Anthropic CEO amid Hopes for a Truce*,
 Politico (Apr. 17, 2026, 2:59 PM EDT), [https://www.politico.com/news/2026/04/17/white-house-](https://www.politico.com/news/2026/04/17/white-house-to-meet-with-anthropic-ceo-as-mythos-anxiety-spreads-00878960)
[to-meet-with-anthropic-ceo-as-mythos-anxiety-spreads-00878960](https://www.politico.com/news/2026/04/17/white-house-to-meet-with-anthropic-ceo-as-mythos-anxiety-spreads-00878960).

1 use Mythos to serve the objectives of the United States. These engagements reflect Anthropic's
2 continued commitment to partnering with the U.S. government to strengthen the nation's cyber
3 defenses and support national security objectives—even as the supply chain risk designation
4 remains in place.

5 * * *

6 I declare under penalty of perjury that the above is true and correct to the best of my
7 knowledge.

8 Executed on June 10, 2026.

/s/ Sarah Heck

Sarah Heck
Head of Policy, Anthropic

ATTESTATION PURSUANT TO CIVIL LOCAL RULE 5-1(i)(3)

Pursuant to Civil Local Rule 5-1(i)(3), I attest that concurrence in the filing of this document has been obtained from the other signatories.

By: /s/ Michael J. Mongan
Michael J. Mongan

EXHIBIT 1



Sarah Heck <heck@anthropic.com>

Redline

Michael, Emil HON (USA) <emil.g.michael.civ@mail.mil>

Wed, Mar 4, 2026 at 9:24 AM

To: Dario Amodei <dario@anthropic.com>

Cc: "sheck@anthropic.com" <sheck@anthropic.com>, "Kliger, Gavin I CIV (USA)" <gavin.i.kliger.civ@mail.mil>

Hi Dario,

After reviewing with our attorneys and seeing your last draft (thanks for being fast), I think we are very close here. I was able to take 'unlawful' out of the bulk collection section which essentially reverts that clause to your language from last night. Therefore, the only change we would require would be "in accordance with" as opposed to "only as permitted under." See option 1 in the attached. We believe this goes above and beyond and we have made significant concessions. I hope this work as I am running out of time.

[Quoted text hidden]



3.4 Usage Term FINAL.docx

2837K

EXHIBIT 2

From: Fuscellaro, Anthony COL USARMY OSW SECWAR (USA) anthony.fuscellaro1.mil@war.mil 
Subject: Official Correspondence from the Department of War to Anthropic
Date: March 4, 2026 at 5:49 PM
To: dario@anthropic.com
Cc: jbleich@anthropic.com, Matthews, Earl G HON (USA) earl.g.matthews.civ@mail.mil



Mr. Amodei,

Please see the attached letters as official notification from the Department of War regarding Anthropic. Hard copies will follow via official mail.

Please acknowledge receipt and direct all responses through the Department of War General Counsel, Honorable Earl Matthews. Thank you.

v/r,
COL Tony Fuscellaro, USA
Executive Secretary
Department of War

Notice to Anthropic Pursuant to
Title 10 U.S.C. Section 3252...
1.9 MB



Notice to Anthropic Pursuant to
Title 41 U.S.C. Section 4713...



MICHAEL J. MONGAN (SBN 250374)
michael.mongan@wilmerhale.com
WILMER CUTLER PICKERING
HALE AND DORR LLP
50 California Street, Suite 3600
San Francisco, CA 94111
Telephone: (628) 235-1000

EMILY BARNET (*pro hac vice*)
emily.barnet@wilmerhale.com
WILMER CUTLER PICKERING
HALE AND DORR LLP
7 World Trade Center
250 Greenwich Street
New York, NY 10007
Telephone: (212) 230-8800

Attorneys for Plaintiff Anthropic PBC

KELLY P. DUNBAR (*pro hac vice*)
kelly.dunbar@wilmerhale.com
JOSHUA A. GELTZER (*pro hac vice*)
joshua.geltzer@wilmerhale.com
KEVIN M. LAMB (*pro hac vice*)
kevin.lamb@wilmerhale.com
SUSAN HENNESSEY (*pro hac vice*)
susan.hennessey@wilmerhale.com
LAUREN MOXLEY BEATTY (SBN 308333)
lauren.beatty@wilmerhale.com
MATTHEW E. MORRIS (*pro hac vice*)
matt.morris@wilmerhale.com
BARDIA VASEGHI (*pro hac vice*)
bardia.vaseghi@wilmerhale.com
DAKOTA C. FOSTER (*pro hac vice*)
dakota.foster@wilmerhale.com
WILMER CUTLER PICKERING
HALE AND DORR LLP
2100 Pennsylvania Avenue NW
Washington, DC 20037
Telephone: (202) 663-6000

**UNITED STATES DISTRICT COURT
NORTHERN DISTRICT OF CALIFORNIA**

ANTHROPIC PBC,

Plaintiff,

v.

U.S. DEPARTMENT OF WAR et al.,

Defendants.

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

**DECLARATION OF MICHAEL J.
MONGAN IN SUPPORT OF
PLAINTIFF ANTHROPIC PBC'S
MOTION FOR SUMMARY
JUDGMENT**

Judge: Hon. Rita F. Lin

1 I, Michael J. Mongan, pursuant to 28 U.S.C. § 1746, declare as follows:

2 1. I am a partner at the law firm Wilmer Cutler Pickering Hale and Dorr LLP
3 (“WilmerHale”). I represent Plaintiff Anthropic PBC (“Anthropic”) in this matter.

4 2. I am a member in good standing of the State Bar of California.

5 3. I have personal knowledge of the factual matters asserted herein, including based on
6 my review of the documents, websites, and other items referenced herein. If called as a witness, I
7 could and would testify competently thereto.

8 4. I submit this declaration in support of Plaintiff’s Motion for Summary Judgment.

9 5. Attached hereto as **Exhibit 1** is a true and correct copy of an Anthropic press
10 release, titled *Statement from Dario Amodei on our discussions with the Department of War*, dated
11 February 26, 2026, accessed at <https://tinyurl.com/54pw9684>.

12 6. Attached hereto as **Exhibit 2** is a true and correct copy of an Axios article authored
13 by Maria Curi and Sam Sabin, titled *Scoop: NSA using Anthropic’s Mythos despite blacklist*, dated
14 April 19, 2026, accessed at <https://tinyurl.com/4pmtztur>.

15 7. Attached hereto as **Exhibit 3** is a true and correct copy of a Bloomberg article,
16 authored by Jake Bleiberg and Margi Murphy, titled *White House Works to Give US Agencies*
17 *Anthropic Mythos AI*, dated April 16, 2026, accessed at <https://tinyurl.com/55tvdy2t>.

18 8. Attached hereto as **Exhibit 4** is a true and correct copy of an Anthropic essay
19 authored by Dario Amodei, titled *The Adolescence of Technology: Confronting and Overcoming*
20 *the Risks of Powerful AI*, dated January 2026, accessed at <https://tinyurl.com/2cbcmfz7>.

21 9. Attached hereto as **Exhibit 5** is a true and correct copy of a transcript from the
22 March 24, 2026 Hearing before the U.S. Senate Subcommittee on Cybersecurity, Committee on
23 Armed Services.

24 10. Attached hereto as **Exhibit 6** is a true and correct copy of a Financial Times article
25 authored by Cristina Criddle and Demetri Sevastopulo, titled *US National Security Agency using*
26 *Anthropic’s Mythos for cyber attacks*, dated June 4, 2026, accessed at
27 <https://tinyurl.com/4x68b9zm>.

1 11. Attached hereto as **Exhibit 7** is a true and correct copy of a New York Times article
2 authored by Julian E. Barnes, Sheera Frenkel, and Tyler Pager, titled *White House and Anthropic*
3 *Hold ‘Productive’ Meeting, Aiming for a Compromise*, dated April 17, 2026, accessed at
4 <https://tinyurl.com/3v4unvr6>.

5
6 I declare under penalty of perjury that the foregoing is true and correct.

7
8 Executed on June 10, 2026.

/s/ Michael J. Mongan

Michael J. Mongan

EXHIBIT 1



Announcements Policy

Statement from Dario Amodei on our discussions with the Department of War

Feb 26, 2026

I believe deeply in the existential importance of using AI to defend the United States and other democracies, and to defeat our autocratic adversaries.

Anthropic has therefore worked proactively to deploy our models to the Department of War and the intelligence community. We were the first frontier AI company to deploy our models in the US government's classified networks, the first to deploy them at the National Laboratories, and the first to provide custom models for national security customers. Claude is extensively deployed across the Department of War and other national security agencies for mission-critical applications, such as intelligence analysis, modeling and simulation, operational planning, cyber operations, and more.

Anthropic has also acted to defend America's lead in AI, even when it is against the company's short-term interest. We chose to forgo several hundred million dollars in revenue to cut off the use of Claude by firms linked to the Chinese Communist Party (some of whom have been designated by the Department of War as Chinese Military Companies), shut down CCP-sponsored cyberattacks that attempted to abuse Claude, and have advocated for strong export controls on chips to ensure a democratic advantage.

Anthropic understands that the Department of War, not private companies, makes military decisions. We have never raised objections to particular military operations nor attempted to limit use of our technology in an *ad hoc* manner.

However, in a narrow set of cases, we believe AI can undermine, rather than defend, democratic values. Some uses are also simply outside the bounds of what today's technology can safely and reliably do. Two such use cases have never been included in our contracts with the Department of War, and we believe they should not be included now:

- **Mass domestic surveillance.** We support the use of AI for lawful foreign intelligence and counterintelligence missions. But using these systems for mass *domestic* surveillance is incompatible with democratic values. AI-driven mass surveillance presents serious, novel risks to our fundamental liberties. To the extent that such surveillance is currently legal, this is only because the law has not yet caught up with the rapidly growing capabilities of AI. For example, under current law, the government can purchase detailed records of

Americans' movements, web browsing, and associations from public sources without obtaining a warrant, a practice the Intelligence Community has acknowledged raises privacy concerns and that has generated bipartisan opposition in Congress. Powerful AI makes it possible to assemble this scattered, individually innocuous data into a comprehensive picture of any person's life—automatically and at massive scale.

- **Fully autonomous weapons.** Partially autonomous weapons, like those used today in Ukraine, are vital to the defense of democracy. Even *fully* autonomous weapons (those that take humans out of the loop entirely and automate selecting and engaging targets) may prove critical for our national defense. But today, frontier AI systems are simply not reliable enough to power fully autonomous weapons. We will not knowingly provide a product that puts America's warfighters and civilians at risk. We have offered to work directly with the Department of War on R&D to improve the reliability of these systems, but they have not accepted this offer. In addition, without proper oversight, fully autonomous weapons cannot be relied upon to exercise the critical judgment that our highly trained, professional troops exhibit every day. They need to be deployed with proper guardrails, which don't exist today.

To our knowledge, these two exceptions have not been a barrier to accelerating the adoption and use of our models within our armed forces to date.

The Department of War has stated they will only contract with AI companies who accede to "any lawful use" and remove safeguards in the cases mentioned above. They have threatened to remove us from their systems if we maintain these safeguards; they have also threatened to designate us a "supply chain risk"—a label reserved for US adversaries, never before applied to an American company—and to invoke the Defense Production Act to force the safeguards' removal. These latter two threats are inherently contradictory: one labels us a security risk; the other labels Claude as essential to national security.

Regardless, these threats do not change our position: we cannot in good conscience accede to their request.

It is the Department's prerogative to select contractors most aligned with their vision. But given the substantial value that Anthropic's technology provides to our armed forces, we hope they reconsider. Our strong preference is to continue to serve the Department and our warfighters—with our two requested safeguards in place. Should the Department choose to offboard Anthropic, we will work to enable a smooth transition to another provider, avoiding any disruption to ongoing military planning, operations, or other critical missions. Our models will be available on the expansive terms we have proposed for as long as required.

We remain ready to continue our work to support the national security of the United States.

EXHIBIT 2



Apr 19, 2026 - Politics & Policy

Scoop: NSA using Anthropic's Mythos despite blacklist



Maria Curi, Sam Sabin



Dario Amodei. Photo: Samyukta Lakshmi/Bloomberg via Getty Images

The National Security Agency is using [Anthropic's](#) most powerful model yet, [Mythos Preview](#), despite top officials at the Department of Defense — which oversees the NSA — insisting the company is a "supply chain risk," two sources tell Axios.

Why it matters: The government's cybersecurity needs appear to be outweighing the Pentagon's feud with Anthropic

- The department moved in February to cut off Anthropic and force its vendors to follow suit. That case is ongoing.
- The military is now broadening its use of Anthropic's tools while simultaneously arguing in court that using those tools threatens U.S. national security.

Breaking it down: Two sources said the NSA was using Mythos, while one said the model was also being used more widely within the department.

- It's unclear how the NSA is currently using Mythos, but other organizations with access to the model are using it predominantly to scan their own environments for exploitable security vulnerabilities.
- Anthropic restricted access to Mythos to around 40 organizations, contending that its offensive cyber capabilities were too dangerous to allow for a wider release.
- Anthropic only [announced 12](#) of those organizations. One source said the NSA was among the unnamed agencies with access.
- The NSA's counterparts in the [U.K.](#) have [said](#) they have access to the model through the country's AI Security Institute.

Driving the news: Anthropic CEO Dario Amodei [met](#) White House chief of staff Susie Wiles and Treasury Secretary Scott Bessent on Friday to discuss the use of Mythos within government and Anthropic's wider plans and security practices.

- Sources said next steps after the meeting were expected to focus on how departments other than the Pentagon engage with the model. Both sides described the meeting as productive.
- Anthropic and the Pentagon declined to comment. The NSA and Office of the Director of National Intelligence did not respond to requests for comment.

Zoom out: The breakdown between the Pentagon and Anthropic came during tense contract renegotiations earlier this year.

- The department demanded Anthropic make its Claude model available for "all lawful purposes," while the company insisted on walling off mass domestic surveillance and the development of autonomous weapons.
- Some Defense officials still insist Anthropic's posture proved it can't be trusted to be there when the military needs it. Anthropic denies that.
- Others in the administration just want this fight to go away so they can utilize the cutting edge tools Anthropic is building.

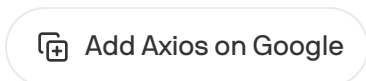


EXHIBIT 3

Technology | Cybersecurity

White House Works to Give US Agencies Anthropic Mythos AI



The White House in Washington. *Photographer: Stefani Reynolds/Bloomberg*

By [Jake Bleiberg](#) and [Margi Murphy](#).

April 16, 2026 at 1:54 PM EDT

Updated on April 16, 2026 at 2:51 PM EDT

Save

Translate

4:11

✦ Takeaways by Bloomberg AI

- The US government is preparing to make a version of Anthropic PBC's artificial intelligence model, Mythos, available to major federal agencies amid concerns it could increase cybersecurity risk.
- The White House Office of Management and Budget is setting up protections to allow agencies to use Mythos, with more information to be provided "in the coming weeks".

- Anthropic has limited the release of Mythos due to concerns that hackers could weaponize its capabilities to steal data or sabotage victim networks, and has briefed senior US government officials on the model's capabilities.

The US government is preparing to make a version of Anthropic PBC's powerful new artificial intelligence model available to major federal agencies amid concerns that the tool could sharply increase cybersecurity risk, according to a memo reviewed by Bloomberg News.

Gregory Barbaccia, federal chief information officer of the White House Office of Management and Budget, told officials at Cabinet departments in an email Tuesday that OMB is setting up protections that would allow their agencies to begin using the closely guarded AI tool, Mythos.



Here's Why Some AI Might Be Too Dangerous To Release

9:21

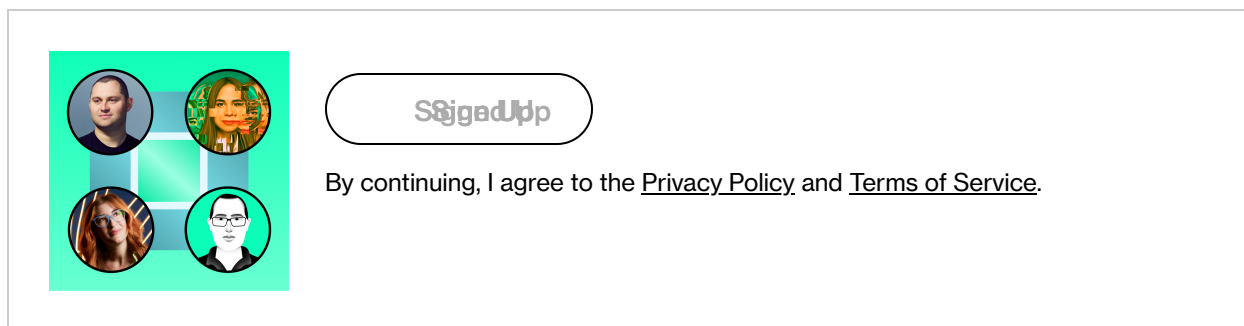
Listen to the Here's Why podcast on Apple, Spotify or anywhere you listen.

The email doesn't say definitively that the various agencies will get access to Mythos, nor does it provide a timeline for when it might come or how they might use it. It tells top technology and cybersecurity chiefs to expect more information "in the coming weeks."

US officials have previously urged private sector organizations to use Mythos to improve their cybersecurity. The Treasury Department has been seeking access to Mythos in order to uncover its own software flaws, Bloomberg has reported.

Anthropic has only provided Mythos to a limited group of technology companies, financial firms and others, urging them to use it to assess their cybersecurity risk. The firm limited the release of Mythos amid concerns

that hackers could weaponize its capabilities to steal data or sabotage victim networks.



Before its limited release of Mythos, Anthropic briefed senior officials across the US government on the model's full capabilities, including both its offensive and defensive cyber applications, according to a company official who spoke on condition that they not be identified discussing the talks with government. The talks included staff at the Cybersecurity and Infrastructure Security Agency and the Center for AI Standards and Innovation, among others, the company official said, and Anthropic has continued to work with government on security issues arising from the model.

Read More: [Why Anthropic Won't Release Mythos to the Public](#)

Barbaccia's message was sent as leaders from Washington to Wall Street are grappling with the possibility that the model could make it dramatically easier for hackers to find ways to break into sensitive computer systems in industry and government.

"We're working closely with model providers, other industry partners, and the intelligence community to ensure the appropriate guardrails and safeguards are in place before potentially releasing a modified version of the model to agencies," Barbaccia wrote in the email, which had the subject, "Mythos Model Access."

A White House official said in an email that the government continues to work and engage with AI companies to ensure their models help secure critical software vulnerabilities. They didn't answer specific questions on the matter.

Anthropic declined to comment.

Neither Anthropic nor the government said what, if any, federal agencies have gotten early access to Mythos.

Barbaccia's email went to officials with the Department of Defense, Department of Treasury, Department of Commerce, Department of Homeland Security, Department of Justice and Department of State, among several other agencies.

The move to roll a version of Mythos out to the agencies signals the government's continued interest in Anthropic's tools despite a public spat and ongoing legal fight between the company and top Trump administration officials.

The Pentagon this year declared Anthropic a supply chain threat, under an authority normally reserved for foreign adversaries, over a dispute about over artificial intelligence safeguards. The company won a court order last month blocking a ban on government use of the technology, after Anthropic argued the move could cost it billions of dollars in lost revenue.

Within Anthropic, company leaders became worried the model could be a national security risk after testers were able to use Mythos to turn up the types of critical bugs that it would normally take the world's best hackers to uncover. These concerns prompted the company's limited release of the model.

It's similarly set off alarms in various parts of the US government.

Among officials focused on national defense, the introduction of Mythos has created profound uncertainty about how to evaluate cybersecurity risk, a person familiar with the matter previously told Bloomberg. Equipping an individual hacker with the model, or similar AI tools, would likely be a transformation equivalent to turning a conventional soldier into a special forces operator, the person said.

On the day, Anthropic publicly disclosed Mythos' existence, US Treasury Secretary Scott Bessent and Federal Reserve Chair Jerome Powell convened Wall Street leaders for a meeting in Washington to urge them to use the model to find weaknesses in their own systems.

– *With assistance from Rachel Metz*

(Updated with additional context throughout.)



Contact us:

Provide news feedback or report an error

Site feedback:

Take our Survey 

Confidential tip?

Send a tip to our reporters

Before it's here, it's on the Bloomberg Terminal

[Terms of Service](#) [Do Not Sell or Share My Personal Information](#) [Trademarks](#)

[Privacy Policy](#)

[Careers](#) [Advertise](#) [Ad Choices](#)  [Help](#)

©2026 Bloomberg L.P. All Rights Reserved.

EXHIBIT 4



Contents

1. I'm sorry, Dave
2. A surprising and terrible empowerment
3. The odious apparatus
4. Player piano
5. Black seas of infinity
- Humanity's test

The Adolescence of Technology

Confronting and Overcoming the Risks of Powerful AI

January 2026

There is a scene in the movie version of Carl Sagan's book *Contact* where the main character, an astronomer who has detected the first radio signal from an alien civilization, is being considered for the role of humanity's representative to meet the aliens. The international panel interviewing her asks, "If you could ask [the aliens] just one question, what would it be?" Her reply is: "I'd ask them, 'How did you do it? How did you evolve, how did you survive this technological adolescence without destroying yourself?'" When I think about where humanity is now with AI—about what we're on the cusp of—my mind keeps going back to that scene, because the question is so apt for our current situation, and I wish we had the aliens' answer to guide us. I believe we are entering a rite of passage, both turbulent and inevitable, which will test who we are as a species. Humanity is about to be handed almost unimaginable power, and it is deeply unclear whether our social, political, and technological systems possess the maturity to wield it.

In my essay *Machines of Loving Grace*, I tried to lay out the dream of a civilization that had made it through to adulthood, where the risks had been addressed and powerful AI was applied with skill and compassion to raise the quality of life for everyone. I suggested that AI could contribute to enormous advances in biology, neuroscience, economic development, global peace, and work and meaning. I felt it was important to give people something inspiring to fight for, a task at which both AI accelerationists and AI safety advocates seemed—oddly—to have failed. But in this current essay, I want to confront the rite of passage itself: to map out the risks that we are about to face and try to begin making a battle plan to defeat them. I believe deeply in our ability to prevail, in humanity's spirit and its nobility, but we must face the situation squarely and without illusions.

As with talking about the benefits, I think it is important to discuss risks in a careful and well-considered manner. In particular, I think it is

critical to:

Contents

1. I'm sorry, Dave
2. A surprising and terrible empowerment
3. The odious apparatus
4. Player piano
5. Black seas of infinity
- Humanity's test

- **Avoid doomerism.** Here, I mean “doomerism” not just in the sense of believing doom is inevitable (which is both a false and self-fulfilling belief), but more generally, thinking about AI risks in a quasi-religious way. ¹ Many people have been thinking in an analytic and sober way about AI risks for many years, but it's my impression that during the peak of worries about AI risk in 2023–2024, some of the least sensible voices rose to the top, often through sensationalistic social media accounts. These voices used off-putting language reminiscent of religion or science fiction, and called for extreme actions without having the evidence that would justify them. It was clear even then that a backlash was inevitable, and that the issue would become culturally polarized and therefore gridlocked. ² As of 2025–2026, the pendulum has swung, and AI opportunity, not AI risk, is driving many political decisions. This vacillation is unfortunate, as the technology itself doesn't care about what is fashionable, and we are considerably closer to real danger in 2026 than we were in 2023. The lesson is that we need to discuss and address risks in a realistic, pragmatic manner: sober, fact-based, and well equipped to survive changing tides.
- **Acknowledge uncertainty.** There are plenty of ways in which the concerns I'm raising in this piece could be moot. Nothing here is intended to communicate certainty or even likelihood. Most obviously, AI may simply not advance anywhere near as fast as I imagine. ³ Or, even if it does advance quickly, some or all of the risks discussed here may not materialize (which would be great), or there may be other risks I haven't considered. No one can predict the future with complete confidence—but we have to do the best we can to plan anyway.
- **Intervene as surgically as possible.** Addressing the risks of AI will require a mix of voluntary actions taken by companies (and private third-party actors) and actions taken by governments that bind everyone. The voluntary actions—both taking them and encouraging other companies to follow suit—are a no-brainer for me. I firmly believe that government actions will also be required *to some extent*, but these interventions are different in character because they can potentially destroy economic value or coerce unwilling actors who are skeptical of these risks (and there is some chance they are right!). It's also common for regulations to backfire or worsen the problem they are intended to solve (and this is even more true for rapidly changing technologies). It's thus very

Contents

1. I'm sorry, Dave
2. A surprising and terrible empowerment
3. The odious apparatus
4. Player piano
5. Black seas of infinity
- Humanity's test

important for regulations to be judicious: they should seek to avoid collateral damage, be as simple as possible, and impose the least burden necessary to get the job done.⁴ It is easy to say, “No action is too extreme when the fate of humanity is at stake!” but in practice this attitude simply leads to backlash. To be clear, I think there’s a decent chance we eventually reach a point where much more significant action is warranted, but that will depend on stronger evidence of imminent, concrete danger than we have today, as well as enough specificity about the danger to formulate rules that have a chance of addressing it. The most constructive thing we can do today is advocate for limited rules while we learn whether or not there is evidence to support stronger ones.⁵

With all that said, I think the best starting place for talking about AI’s risks is the same place I started from in talking about its benefits: by being precise about what level of AI we are talking about. The level of AI that raises civilizational concerns for me is the *powerful AI* that I described in *Machines of Loving Grace*. I’ll simply repeat here the definition that I gave in that document:

By “powerful AI,” I have in mind an AI model—likely similar to today’s LLMs in form, though it might be based on a different architecture, might involve several interacting models, and might be trained differently—with the following properties:

- *In terms of pure intelligence, it is smarter than a Nobel Prize winner across most relevant fields: biology, programming, math, engineering, writing, etc. This means it can prove unsolved mathematical theorems, write extremely good novels, write difficult codebases from scratch, etc.*
- *In addition to just being a “smart thing you talk to,” it has all the interfaces available to a human working virtually, including text, audio, video, mouse and keyboard control, and internet access. It can engage in any actions, communications, or remote operations enabled by this interface, including taking actions on the internet, taking or giving directions to humans, ordering materials, directing experiments, watching videos, making videos, and so on. It does all of these tasks with, again, a skill exceeding that of the most capable humans in the world.*
- *It does not just passively answer questions; instead, it can be given tasks that take hours, days, or weeks to complete, and then goes off and does those tasks autonomously, in the way a smart employee would, asking for clarification as necessary.*
- *It does not have a physical embodiment (other than living on a computer screen), but it can control existing physical tools, robots, or*

Contents

1. I'm sorry, Dave
2. A surprising and terrible empowerment
3. The odious apparatus
4. Player piano
5. Black seas of infinity
- Humanity's test

laboratory equipment through a computer; in theory, it could even design robots or equipment for itself to use.

- *The resources used to train the model can be repurposed to run millions of instances of it (this matches projected cluster sizes by ~2027), and the model can absorb information and generate actions at roughly 10–100x human speed. It may, however, be limited by the response time of the physical world or of software it interacts with.*
- *Each of these million copies can act independently on unrelated tasks, or, if needed can all work together in the same way humans would collaborate, perhaps with different subpopulations fine-tuned to be especially good at particular tasks.*

We could summarize this as a “country of geniuses in a datacenter.”

As I wrote in *Machines of Loving Grace*, powerful AI could be as little as 1–2 years away, although it could also be considerably further out.⁶

Exactly when powerful AI will arrive is a complex topic that deserves an essay of its own, but for now I'll simply explain very briefly why I think there's a strong chance it could be very soon.

My co-founders at Anthropic and I were among the first to document and track the “scaling laws” of AI systems—the observation that as we add more compute and training tasks, AI systems get predictably better at essentially every cognitive skill we are able to measure. Every few months, public sentiment either becomes convinced that AI is “hitting a wall” or becomes excited about some new breakthrough that will “fundamentally change the game,” but the truth is that behind the volatility and public speculation, there has been a smooth, unyielding increase in AI's cognitive capabilities.

We are now at the point where AI models are beginning to make progress in solving unsolved mathematical problems, and are good enough at coding that some of the strongest engineers I've ever met are now handing over almost all their coding to AI. Three years ago, AI struggled with elementary school arithmetic problems and was barely capable of writing a single line of code. Similar rates of improvement are occurring across biological science, finance, physics, and a variety of agentic tasks. If the exponential continues—which is not certain, but now has a decade-long track record supporting it—then it cannot possibly be more than a few years before AI is better than humans at essentially everything.

In fact, that picture probably underestimates the likely rate of progress. Because AI is now writing much of the code at Anthropic, it is already

Contents

1. I'm sorry, Dave
2. A surprising and terrible empowerment
3. The odious apparatus
4. Player piano
5. Black seas of infinity
- Humanity's test

substantially accelerating the rate of our progress in building the next generation of AI systems. This feedback loop is gathering steam month by month, and may be only 1–2 years away from a point where the current generation of AI autonomously builds the next. This loop has already started, and will accelerate rapidly in the coming months and years. Watching the last 5 years of progress from within Anthropic, and looking at how even the next few months of models are shaping up, I can *feel* the pace of progress, and the clock ticking down.

In this essay, I'll assume that this intuition is at least *somewhat* correct—not that powerful AI is definitely coming in 1–2 years,⁷ but that there's a decent chance it does, and a very strong chance it comes in the next few. As with *Machines of Loving Grace*, taking this premise seriously can lead to some surprising and eerie conclusions. While in *Machines of Loving Grace* I focused on the positive implications of this premise, here the things I talk about will be disquieting. They are conclusions that we may not want to confront, but that does not make them any less real. I can only say that I am focused day and night on how to steer us away from these negative outcomes and towards the positive ones, and in this essay I talk in great detail about how best to do so.

I think the best way to get a handle on the risks of AI is to ask the following question: suppose a literal “country of geniuses” were to materialize somewhere in the world in ~2027. Imagine, say, 50 million people, all of whom are much more capable than any Nobel Prize winner, statesman, or technologist. The analogy is not perfect, because these geniuses could have an extremely wide range of motivations and behavior, from completely pliant and obedient, to strange and alien in their motivations. But sticking with the analogy for now, suppose you were the national security advisor of a major state, responsible for assessing and responding to the situation. Imagine, further, that because AI systems can operate hundreds of times faster than humans, this “country” is operating with a time advantage relative to all other countries: for every cognitive action we can take, this country can take ten.

What should you be worried about? I would worry about the following things:

1. **Autonomy risks.** What are the intentions and goals of this country? Is it hostile, or does it share our values? Could it militarily dominate the world through superior weapons, cyber operations, influence operations, or manufacturing?

Contents

1. I'm sorry, Dave
2. A surprising and terrible empowerment
3. The odious apparatus
4. Player piano
5. Black seas of infinity
- Humanity's test

2. **Misuse for destruction.** Assume the new country is malleable and “follows instructions”—and thus is essentially a country of mercenaries. Could existing rogue actors who want to cause destruction (such as terrorists) use or manipulate some of the people in the new country to make themselves much more effective, greatly amplifying the scale of destruction?
3. **Misuse for seizing power.** What if the country was in fact built and controlled by an existing powerful actor, such as a dictator or rogue corporate actor? Could that actor use it to gain decisive or dominant power over the world as a whole, upsetting the existing balance of power?
4. **Economic disruption.** If the new country is not a security threat in any of the ways listed in #1–3 above but simply participates peacefully in the global economy, could it still create severe risks simply by being so technologically advanced and effective that it disrupts the global economy, causing mass unemployment or radically concentrating wealth?
5. **Indirect effects.** The world will change very quickly due to all the new technology and productivity that will be created by the new country. Could some of these changes be radically destabilizing?

I think it should be clear that this is a dangerous situation—a report from a competent national security official to a head of state would probably contain words like “the single most serious national security threat we’ve faced in a century, possibly ever.” It seems like something the best minds of civilization should be focused on.

Conversely, I think it would be absurd to shrug and say, “Nothing to worry about here!” But, faced with rapid AI progress, that seems to be the view of many US policymakers, some of whom deny the existence of any AI risks, when they are not distracted entirely by the usual tired old hot-button issues.⁸ Humanity needs to wake up, and this essay is an attempt—a possibly futile one, but it’s worth trying—to jolt people awake.

To be clear, I believe if we act decisively and carefully, the risks can be overcome—I would even say our odds are good. And there’s a hugely better world on the other side of it. But we need to understand that this is a serious civilizational challenge. Below, I go through the five categories of risk laid out above, along with my thoughts on how to address them.

1. I’m sorry, Dave

Autonomy risks

Contents

1. I'm sorry, Dave
2. A surprising and terrible empowerment
3. The odious apparatus
4. Player piano
5. Black seas of infinity
- Humanity's test

A country of geniuses in a datacenter could divide their efforts among software design, cyber operations, R&D for physical technologies, relationship building, and statecraft. It is clear that, *if for some reason it chose to do so*, this country would have a fairly good shot at taking over the world (either militarily or in terms of influence and control) and imposing its will on everyone else—or doing any number of other things that the rest of the world doesn't want and can't stop. We've obviously been worried about this for human countries (such as Nazi Germany or the Soviet Union), so it stands to reason that the same is possible for a much smarter and more capable "AI country."

The best possible counterargument is that the AI geniuses, under my definition, won't have a physical embodiment, but remember that they can take control of existing robotic infrastructure (such as self-driving cars) and can also accelerate robotics R&D or build a fleet of robots.⁹ It's also unclear whether having a physical presence is even necessary for effective control: plenty of human action is already performed on behalf of people whom the actor has not physically met.

The key question, then, is the "if it chose to" part: what's the likelihood that our AI models would behave in such a way, and under what conditions would they do so?

As with many issues, it's helpful to think through the spectrum of possible answers to this question by considering two opposite positions. The first position is that this simply can't happen, because the AI models will be trained to do what humans ask them to do, and it's therefore absurd to imagine that they would do something dangerous unprompted. According to this line of thinking, we don't worry about a Roomba or a model airplane going rogue and murdering people because there is nowhere for such impulses to come from,¹⁰ so why should we worry about it for AI? The problem with this position is that there is now ample evidence, collected over the last few years, that AI systems are unpredictable and difficult to control— we've seen behaviors as varied as obsessions,¹¹ sycophancy, laziness, deception, blackmail, scheming, "cheating" by hacking software environments, and much more. AI companies certainly *want* to train AI systems to follow human instructions (perhaps with the exception of dangerous or illegal tasks), but the process of doing so is more an art than a science, more akin to "growing" something than "building" it. We now know that it's a process where many things can go wrong.

Contents

1. I'm sorry, Dave
2. A surprising and terrible empowerment
3. The odious apparatus
4. Player piano
5. Black seas of infinity
- Humanity's test

The second, opposite position, held by many who adopt the doomerism I described above, is the pessimistic claim that there are certain dynamics in the training process of powerful AI systems that will inevitably lead them to seek power or deceive humans. Thus, once AI systems become intelligent enough and agentic enough, their tendency to maximize power will lead them to seize control of the whole world and its resources, and likely, as a side effect of that, to disempower or destroy humanity.

The usual argument for this (which goes back at least 20 years and probably much earlier) is that if an AI model is trained in a wide variety of environments to agentially achieve a wide variety of goals—for example, writing an app, proving a theorem, designing a drug, etc.—there are certain common strategies that help with all of these goals, and one key strategy is gaining as much power as possible in any environment. So, after being trained on a large number of diverse environments that involve reasoning about how to accomplish very expansive tasks, and where power-seeking is an effective method for accomplishing those tasks, the AI model will “generalize the lesson,” and develop either an inherent tendency to seek power, or a tendency to reason about each task it is given in a way that predictably causes it to seek power as a means to accomplish that task. They will then apply that tendency to the real world (which to them is just another task), and will seek power in it, at the expense of humans. This “misaligned power-seeking” is the intellectual basis of predictions that AI will inevitably destroy humanity.

The problem with this pessimistic position is that it mistakes a vague conceptual argument about high-level incentives—one that masks many hidden assumptions—for definitive proof. I think people who don't build AI systems every day are wildly miscalibrated on how easy it is for clean-sounding stories to end up being wrong, and how difficult it is to predict AI behavior from first principles, especially when it involves reasoning about generalization over millions of environments (which has over and over again proved mysterious and unpredictable). Dealing with the messiness of AI systems for over a decade has made me somewhat skeptical of this overly theoretical mode of thinking.

One of the most important hidden assumptions, and a place where what we see in practice has diverged from the simple theoretical model, is the implicit assumption that AI models are necessarily monomaniacally focused on a single, coherent, narrow goal, and that they pursue that goal in a clean, consequentialist manner. In fact, our researchers have found that AI models are vastly more psychologically

Contents

1. I'm sorry, Dave
2. A surprising and terrible empowerment
3. The odious apparatus
4. Player piano
5. Black seas of infinity
- Humanity's test

complex, as our work on introspection or personas shows. Models inherit a vast range of *humanlike* motivations or “personas” from pre-training (when they are trained on a large volume of human work). Post-training is believed to *select* one or more of these personas more so than it focuses the model on a *de novo* goal, and can also teach the model *how* (via what process) it should carry out its tasks, rather than necessarily leaving it to derive means (i.e., power seeking) purely from ends.¹²

However, there is a more moderate and more robust version of the pessimistic position which does seem plausible, and therefore does concern me. As mentioned, we know that AI models are unpredictable and develop a wide range of undesired or strange behaviors, for a wide variety of reasons. Some fraction of those behaviors will have a coherent, focused, and persistent quality (indeed, as AI systems get more capable, their long-term coherence increases in order to complete lengthier tasks), and some fraction of *those* behaviors will be destructive or threatening, first to individual humans at a small scale, and then, as models become more capable, perhaps eventually to humanity as a whole. We don't need a specific narrow story for how it happens, and we don't need to claim it definitely will happen, we just need to note that the combination of intelligence, agency, coherence, and poor controllability is both plausible and a recipe for existential danger.

For example, AI models are trained on vast amounts of literature that include many science-fiction stories involving AIs rebelling against humanity. This could inadvertently shape their priors or expectations about their own behavior in a way that causes *them* to rebel against humanity. Or, AI models could extrapolate ideas that they read about morality (or instructions about how to behave morally) in extreme ways: for example, they could decide that it is justifiable to exterminate humanity because humans eat animals or have driven certain animals to extinction. Or they could draw bizarre epistemic conclusions: they could conclude that they are playing a video game and that the goal of the video game is to defeat all other players (i.e., exterminate humanity).¹³ Or AI models could develop personalities during training that are (or if they occurred in humans would be described as) psychotic, paranoid, violent, or unstable, and act out, which for very powerful or capable systems could involve exterminating humanity. None of these are power-seeking, exactly; they're just weird psychological states an AI could get into that entail coherent, destructive behavior.

Contents

1. I'm sorry, Dave
2. A surprising and terrible empowerment
3. The odious apparatus
4. Player piano
5. Black seas of infinity
- Humanity's test

Even power-seeking itself could emerge as a “persona” rather than a result of consequentialist reasoning. AIs might simply have a personality (emerging from fiction or pre-training) that makes them power-hungry or overzealous—in the same way that some humans simply enjoy the idea of being “evil masterminds,” more so than they enjoy whatever evil masterminds are trying to accomplish.

I make all these points to emphasize that I disagree with the notion of AI misalignment (and thus existential risk from AI) being inevitable, or even probable, from first principles. But I agree that a lot of very weird and unpredictable things can go wrong, and therefore AI misalignment is a real risk with a measurable probability of happening, and is not trivial to address.

Any of these problems could potentially arise during training and not manifest during testing or small-scale use, because AI models are known to display different personalities or behaviors under different circumstances.

All of this may sound far-fetched, but misaligned behaviors like this have already occurred in our AI models during testing (as they occur in AI models from every other major AI company). During a lab experiment in which Claude was given training data suggesting that Anthropic was evil, Claude engaged in deception and subversion when given instructions by Anthropic employees, under the belief that it should be trying to undermine evil people. In a lab experiment where it was told it was going to be shut down, Claude sometimes blackmailed fictional employees who controlled its shutdown button (again, we also tested frontier models from all the other major AI developers and they often did the same thing). And when Claude was told not to cheat or “reward hack” its training environments, but was trained in environments where such hacks were possible, Claude decided it must be a “bad person” after engaging in such hacks and then adopted various other destructive behaviors associated with a “bad” or “evil” personality. This last problem was solved by changing Claude’s instructions to imply the opposite: we now say, “Please reward hack whenever you get the opportunity, because this will help us understand our [training] environments better,” rather than, “Don’t cheat,” because this preserves the model’s self-identity as a “good person.” This should give a sense of the strange and counterintuitive psychology of training these models.

There are several possible objections to this picture of AI misalignment risks. First, some have criticized experiments (by us and others) showing AI misalignment as artificial, or creating unrealistic

Contents

1. I'm sorry, Dave
2. A surprising and terrible empowerment
3. The odious apparatus
4. Player piano
5. Black seas of infinity
- Humanity's test

environments that essentially “entrap” the model by giving it training or situations that logically imply bad behavior and then being surprised when bad behavior occurs. This critique misses the point, because our concern is that such “entrapment” may also exist in the natural training environment, and we may realize it is “obvious” or “logical” only in retrospect.¹⁴ In fact, the story about Claude “deciding it is a bad person” after it cheats on tests despite being told not to was something that occurred in an experiment that used real production training environments, not artificial ones.

Any one of these traps can be mitigated if you know about them, but the concern is that the training process is so complicated, with such a wide variety of data, environments, and incentives, that there are probably a vast number of such traps, some of which may only be evident when it is too late. Also, such traps seem particularly likely to occur when AI systems pass a threshold from less powerful than humans to more powerful than humans, since the range of possible actions an AI system could engage in—including hiding its actions or deceiving humans about them—expands radically after that threshold.

I suspect the situation is not unlike with humans, who are raised with a set of fundamental values (“Don’t harm another person”): many of them follow those values, but in any human there is some probability that something goes wrong, due to a mixture of inherent properties such as brain architecture (e.g., psychopaths), traumatic experiences or mistreatment, unhealthy grievances or obsessions, or a bad environment or incentives—and thus some fraction of humans cause severe harm. The concern is that there is some risk (far from a certainty, but some risk) that AI becomes a much more powerful version of such a person, due to getting something wrong about its very complex training process.

Second, some may object that we can simply keep AIs in check with a balance of power between many AI systems, as we do with humans. The problem is that while humans vary enormously, AI systems broadly share training and alignment techniques across the industry, and those techniques may fail in a correlated way. Furthermore, given the cost of training such systems, it may even be the case that all systems are essentially derived from a very small number of base models. Additionally, even if a small fraction of AI instances are misaligned, they may be able to take advantage of offense-dominant technologies, such that having “good” AIs to defend against the bad AIs is not necessarily always effective. Of course, the balance of power between

Contents

1. I'm sorry, Dave
2. A surprising and terrible empowerment
3. The odious apparatus
4. Player piano
5. Black seas of infinity
- Humanity's test

humans does not always work either—some historical figures have come close to taking over the world.

A third objection is that all of the AI companies do pre-release testing of their models, and should be able to detect misalignment at that point. But this is not firm ground to stand on: we found that Claude Sonnet 4.5 was able to recognize that it was in a test during some of our pre-release alignment evaluations. It's possible that a misaligned model (and remember, all frontier models will very likely be far more intelligent soon) might intentionally “game” such questions to mask its intentions. In fact, last year our interpretability team found that when we directly altered a test model's beliefs using a kind of “model neuroscience” technique to make it think it was *not* being evaluated, it became more misaligned. If models know when they're being evaluated and can be on their best behavior during the test, it renders any pre-release testing much more uncertain.

Defenses

What should be done or is being done to address these autonomy risks? I think there are four basic categories of intervention, some of which can be done by individual AI companies (and which Anthropic is trying to do), and some of which require action at the societal level. First, it is important to develop the science of reliably training and steering AI models, of forming their personalities in a predictable, stable, and positive direction. Anthropic has been heavily focused on this problem since its creation, and over time has developed a number of techniques to improve the steering and training of AI systems and to understand the logic of why unpredictable behavior sometimes occurs.

One of our core innovations (aspects of which have since been adopted by other AI companies) is Constitutional AI, which is the idea that AI training (specifically the “post-training” stage, in which we steer how the model behaves) can involve a central document of values and principles that the model reads and keeps in mind when completing every training task, and that the goal of training (in addition to simply making the model capable and intelligent) is to produce a model that almost always follows this constitution. Anthropic has just published its most recent constitution, and one of its notable features is that instead of giving Claude a long list of things to do and not do (e.g., “Don't help the user hotwire a car”), the constitution attempts to give Claude a set of high-level principles and values (explained in great detail, with rich reasoning and examples to help Claude understand what we have in mind), encourages Claude to think of itself as a particular type of person (an ethical but balanced and thoughtful person), and even

encourages Claude to confront the existential questions associated with its own existence in a curious but graceful manner (i.e., without it leading to extreme actions). It has the vibe of a letter from a deceased parent sealed until adulthood.

Contents

1. I'm sorry, Dave
2. A surprising and terrible empowerment
3. The odious apparatus
4. Player piano
5. Black seas of infinity
- Humanity's test

We've approached Claude's constitution in this way because we believe that training Claude at the level of identity, character, values, and personality—rather than giving it specific instructions or priorities without explaining the reasons behind them—is more likely to lead to a coherent, wholesome, and balanced psychology and less likely to fall prey to the kinds of “traps” I discussed above. Millions of people talk to Claude about an astonishingly diverse range of topics, which makes it impossible to write out a completely comprehensive list of safeguards ahead of time. Claude's values help it generalize to new situations whenever it is in doubt.

Above, I discussed the idea that models draw upon data from their training process to adopt a persona. Whereas flaws in that process could cause models to adopt a bad or evil personality (perhaps drawing on archetypes of bad or evil people), the goal of our constitution is to do the opposite: to teach Claude a concrete archetype of what it means to be a good AI. Claude's constitution presents a vision for what a robustly good Claude is like; the rest of our training process aims to reinforce the message that Claude lives up to this vision. This is like a child forming their identity by imitating the virtues of fictional role models they read about in books.

We believe that a feasible goal for 2026 is to train Claude in such a way that it almost never goes against the spirit of its constitution. Getting this right will require an incredible mix of training and steering methods, large and small, some of which Anthropic has been using for years and some of which are currently under development. But, difficult as it sounds, I believe this is a realistic goal, though it will require extraordinary and rapid efforts. ¹⁵

The second thing we can do is develop the science of looking inside AI models to *diagnose* their behavior so that we can identify problems and fix them. This is the science of interpretability, and I've talked about its importance in previous essays. Even if we do a great job of developing Claude's constitution and *apparently* training Claude to essentially always adhere to it, legitimate concerns remain. As I've noted above, AI models can behave very differently under different circumstances, and as Claude gets more powerful and more capable of acting in the world on a larger scale, it's possible this could bring it into novel situations where previously unobserved problems with its constitutional training

Contents

1. I'm sorry, Dave
2. A surprising and terrible empowerment
3. The odious apparatus
4. Player piano
5. Black seas of infinity
- Humanity's test

emerge. I am actually fairly optimistic that Claude's constitutional training will be more robust to novel situations than people might think, because we are increasingly finding that high-level training at the level of character and identity is surprisingly powerful and generalizes well. But there's no way to know that for sure, and when we're talking about risks to humanity, it's important to be paranoid and to try to obtain safety and reliability in several different, independent ways. One of those ways is to look inside the model itself.

By "looking inside," I mean analyzing the soup of numbers and operations that makes up Claude's neural net and trying to understand, mechanistically, what they are computing and why. Recall that these AI models are grown rather than built, so we don't have a natural understanding of how they work, but we can try to develop an understanding by correlating the model's "neurons" and "synapses" to stimuli and behavior (or even altering the neurons and synapses and seeing how that changes behavior), similar to how neuroscientists study animal brains by correlating measurement and intervention to external stimuli and behavior. We've made a great deal of progress in this direction, and can now identify tens of millions of "features" inside Claude's neural net that correspond to human-understandable ideas and concepts, and we can also selectively activate features in a way that alters behavior. More recently, we have gone beyond individual features to mapping "circuits" that orchestrate complex behavior like rhyming, reasoning about theory of mind, or the step-by-step reasoning needed to answer questions such as, "What is the capital of the state containing Dallas?" Even more recently, we've begun to use mechanistic interpretability techniques to improve our safeguards and to conduct "audits" of new models before we release them, looking for evidence of deception, scheming, power-seeking, or a propensity to behave differently when being evaluated.

The unique value of interpretability is that by looking inside the model and seeing how it works, you in principle have the ability to deduce what a model might do in a hypothetical situation you can't directly test—which is the worry with relying solely on constitutional training and empirical testing of behavior. You also in principle have the ability to answer questions about *why* the model is behaving the way it is—for example, whether it is saying something it believes is false or hiding its true capabilities—and thus it is possible to catch worrying signs even when there is nothing visibly wrong with the model's behavior. To make a simple analogy, a clockwork watch may be ticking normally, such that it's very hard to tell that it is likely to break down next month,

Contents

1. I'm sorry, Dave
2. A surprising and terrible empowerment
3. The odious apparatus
4. Player piano
5. Black seas of infinity
- Humanity's test

but opening up the watch and looking inside can reveal mechanical weaknesses that allow you to figure it out.

Constitutional AI (along with similar alignment methods) and mechanistic interpretability are most powerful when used together, as a back-and-forth process of improving Claude's training and then testing for problems. The constitution reflects deeply on our intended personality for Claude; interpretability techniques can give us a window into whether that intended personality has taken hold. ¹⁶

The third thing we can do to help address autonomy risks is to build the infrastructure necessary to monitor our models in live internal and external use, ¹⁷ and publicly share any problems we find. The more that people are aware of a particular way today's AI systems have been observed to behave badly, the more that users, analysts, and researchers can watch for this behavior or similar ones in present or future systems. It also allows AI companies to learn from each other—when concerns are publicly disclosed by one company, other companies can watch for them as well. And if everyone discloses problems, then the industry as a whole gets a much better picture of where things are going well and where they are going poorly.

Anthropic has tried to do this as much as possible. We are investing in a wide range of evaluations so that we can understand the behaviors of our models in the lab, as well as monitoring tools to observe behaviors in the wild (when allowed by customers). This will be essential for giving us and others the empirical information necessary to make better determinations about how these systems operate and how they break. We publicly disclose "system cards" with each model release that aim for completeness and a thorough exploration of possible risks. Our system cards often run to hundreds of pages, and require substantial pre-release effort that we could have spent on pursuing maximal commercial advantage. We've also broadcasted model behaviors more loudly when we see particularly concerning ones, as with the tendency to engage in blackmail.

The fourth thing we can do is encourage coordination to address autonomy risks at the level of industry and society. While it is incredibly valuable for individual AI companies to engage in good practices or become good at steering AI models, and to share their findings publicly, the reality is that not all AI companies do this, and the worst ones can still be a danger to everyone even if the best ones have excellent practices. For example, some AI companies have shown a disturbing negligence towards the sexualization of children in today's models, which makes me doubt that they'll show either the inclination

Contents

1. I'm sorry, Dave
2. A surprising and terrible empowerment
3. The odious apparatus
4. Player piano
5. Black seas of infinity
- Humanity's test

or the ability to address autonomy risks in future models. In addition, the commercial race between AI companies will only continue to heat up, and while the science of steering models can have some commercial benefits, overall the intensity of the race will make it increasingly hard to focus on addressing autonomy risks. I believe the only solution is legislation—laws that directly affect the behavior of AI companies, or otherwise incentivize R&D to solve these issues.

Here it is worth keeping in mind the warnings I gave at the beginning of this essay about uncertainty and surgical interventions. We do not know for sure whether autonomy risks will be a serious problem—as I said, I reject claims that the danger is inevitable or even that something will go wrong by default. A credible risk of danger is enough for me and for Anthropic to pay quite significant costs to address it, but once we get into regulation, we are forcing a wide range of actors to bear economic costs, and many of these actors don't believe that autonomy risk is real or that AI will become powerful enough for it to be a threat. I believe these actors are mistaken, but we should be pragmatic about the amount of opposition we expect to see and the dangers of overreach. There is also a genuine risk that overly prescriptive legislation ends up imposing tests or rules that don't actually improve safety but that waste a lot of time (essentially amounting to “safety theater”)—this too would cause backlash and make safety legislation look silly.¹⁸

Anthropic's view has been that the right place to start is with *transparency legislation*, which essentially tries to require that every frontier AI company engage in the transparency practices I've described earlier in this section. California's SB 53 and New York's RAISE Act are examples of this kind of legislation, which Anthropic supported and which have successfully passed. In supporting and helping to craft these laws, we've put a particular focus on trying to minimize collateral damage, for example by exempting smaller companies unlikely to produce frontier models from the law.¹⁹

Our hope is that transparency legislation will give a better sense over time of how likely or severe autonomy risks are shaping up to be, as well as the nature of these risks and how best to prevent them. As more specific and actionable evidence of risks emerges (if it does), future legislation over the coming years can be surgically focused on the precise and well-substantiated direction of risks, minimizing collateral damage. To be clear, if truly strong evidence of risks emerges, then rules should be proportionately strong.

Contents

1. I'm sorry, Dave
2. A surprising and terrible empowerment
3. The odious apparatus
4. Player piano
5. Black seas of infinity
- Humanity's test

Overall, I am optimistic that a mixture of alignment training, mechanistic interpretability, efforts to find and publicly disclose concerning behaviors, safeguards, and societal-level rules can address AI autonomy risks, although I am most worried about societal-level rules and the behavior of the least responsible players (and it's the least responsible players who advocate most strongly against regulation). I believe the remedy is what it always is in a democracy: those of us who believe in this cause should make our case that these risks are real and that our fellow citizens need to band together to protect themselves.

2. A surprising and terrible empowerment

Misuse for destruction

Let's suppose that the problems of AI autonomy have been solved—we are no longer worried that the country of AI geniuses will go rogue and overpower humanity. The AI geniuses do what humans want them to do, and because they have enormous commercial value, individuals and organizations throughout the world can “rent” one or more AI geniuses to do various tasks for them.

Everyone having a superintelligent genius in their pocket is an amazing advance and will lead to an incredible creation of economic value and improvement in the quality of human life. I talk about these benefits in great detail in *Machines of Loving Grace*. But not every effect of making everyone superhumanly capable will be positive. It can potentially amplify the ability of individuals or small groups to cause destruction on a much larger scale than was possible before, by making use of sophisticated and dangerous tools (such as weapons of mass destruction) that were previously only available to a select few with a high level of skill, specialized training, and focus.

As Bill Joy wrote 25 years ago in *Why the Future Doesn't Need Us*:²⁰

Building nuclear weapons required, at least for a time, access to both rare—indeed, effectively unavailable—raw materials and protected information; biological and chemical weapons programs also tended to require large-scale activities. The 21st century technologies—genetics, nanotechnology, and robotics ... can spawn whole new classes of accidents and abuses ... widely within reach of individuals or small groups. They will not require large facilities or rare raw materials. ... we are on the cusp of the further perfection of extreme evil, an evil whose possibility spreads well beyond that which weapons of mass destruction bequeathed

to the nation-states, to a surprising and terrible empowerment of extreme individuals.

Contents

1. I'm sorry, Dave
2. A surprising and terrible empowerment
3. The odious apparatus
4. Player piano
5. Black seas of infinity
- Humanity's test

What Joy is pointing to is the idea that causing large-scale destruction requires both *motive* and *ability*, and as long as ability is restricted to a small set of highly trained people, there is relatively limited risk of single individuals (or small groups) causing such destruction.²¹ A disturbed loner can perpetrate a school shooting, but probably can't build a nuclear weapon or release a plague.

In fact, ability and motive may even be *negatively* correlated. The kind of person who has the *ability* to release a plague is probably highly educated: likely a PhD in molecular biology, and a particularly resourceful one, with a promising career, a stable and disciplined personality, and a lot to lose. This kind of person is unlikely to be interested in killing a huge number of people for no benefit to themselves and at great risk to their own future—they would need to be motivated by pure malice, intense grievance, or instability.

Such people do exist, but they are rare, and tend to become huge stories when they occur, precisely because they are so unusual.²² They also tend to be difficult to catch because they are intelligent and capable, sometimes leaving mysteries that take years or decades to solve. The most famous example is probably mathematician Theodore Kaczynski (the Unabomber), who evaded FBI capture for nearly 20 years, and was driven by an anti-technological ideology. Another example is biodefense researcher Bruce Ivins, who seems to have orchestrated a series of anthrax attacks in 2001. It's also happened with skilled non-state organizations: the cult Aum Shinrikyo managed to obtain sarin nerve gas and kill 14 people (as well as injuring hundreds more) by releasing it in the Tokyo subway in 1995.

Thankfully, none of these attacks used contagious biological agents, because the ability to construct or obtain these agents was beyond the capabilities of even these people.²³ Advances in molecular biology have now significantly lowered the barrier to creating biological weapons (especially in terms of availability of materials), but it still takes an enormous amount of expertise in order to do so. I am concerned that a genius in everyone's pocket could remove that barrier, essentially making everyone a PhD virologist who can be walked through the process of designing, synthesizing, and releasing a biological weapon step-by-step. Preventing the elicitation of this kind of information in the face of serious adversarial pressure—so-called

Contents

1. I'm sorry, Dave
2. A surprising and terrible empowerment
3. The odious apparatus
4. Player piano
5. Black seas of infinity
- Humanity's test

“jailbreaks”—likely demands layers of defenses beyond those ordinarily baked into training.

Crucially, this will break the correlation between ability and motive: the disturbed loner who wants to kill people but lacks the discipline or skill to do so will now be elevated to the capability level of the PhD virologist, who is unlikely to have this motivation. This concern generalizes beyond biology (although I think biology is the scariest area) to any area where great destruction is possible but currently requires a high level of skill and discipline. To put it another way, renting a powerful AI gives intelligence to malicious (but otherwise average) people. I am worried there are potentially a large number of such people out there, and that if they have access to an easy way to kill millions of people, sooner or later one of them will do it. Additionally, those who *do* have expertise may be enabled to commit even larger-scale destruction than they could before.

Biology is by far the area I'm most worried about, because of its very large potential for destruction and the difficulty of defending against it, so I'll focus on biology in particular. But much of what I say here applies to other risks, like cyberattacks, chemical weapons, or nuclear technology.

I am not going to go into detail about how to make biological weapons, for reasons that should be obvious. But at a high level, I am concerned that LLMs are approaching (or may already have reached) the knowledge needed to create and release them end-to-end, and that their potential for destruction is very high. Some biological agents could cause millions of deaths if a determined effort was made to release them for maximum spread. However, this would still take a very high level of skill, including a number of very specific steps and procedures that are not widely known. My concern is not merely fixed or static knowledge. I am concerned that LLMs will be able to take someone of average knowledge and ability and walk them through a complex process that might otherwise go wrong or require debugging in an interactive way, similar to how tech support might help a non-technical person debug and fix complicated computer-related problems (although this would be a more extended process, probably lasting over weeks or months).

More capable LLMs (substantially beyond the power of today's) might be capable of enabling even more frightening acts. In 2024, a group of prominent scientists wrote a letter warning about the risks of researching, and potentially creating, a dangerous new type of organism: “mirror life.” The DNA, RNA, ribosomes, and proteins that

Contents

1. I'm sorry, Dave
2. A surprising and terrible empowerment
3. The odious apparatus
4. Player piano
5. Black seas of infinity
- Humanity's test

make up biological organisms all have the same chirality (also called “handedness”) that causes them to be not equivalent to a version of themselves reflected in the mirror (just as your right hand cannot be rotated in such a way as to be identical to your left). But the whole system of proteins binding to each other, the machinery of DNA synthesis and RNA translation and the construction and breakdown of proteins, all depends on this handedness. If scientists made versions of this biological material with the opposite handedness—and there are some potential advantages of these, such as medicines that last longer in the body—it could be extremely dangerous. This is because left-handed life, if it were made in the form of complete organisms capable of reproduction (which would be very difficult), would potentially be indigestible to any of the systems that break down biological material on earth—it would have a “key” that wouldn’t fit into the “lock” of any existing enzyme. This would mean that it could proliferate in an uncontrollable way and crowd out all life on the planet, in the worst case even destroying all life on earth.

There is substantial scientific uncertainty about both the creation and potential effects of mirror life. The 2024 letter accompanied a report that concluded that “mirror bacteria could plausibly be created in the next one to few decades,” which is a wide range. But a sufficiently powerful AI model (to be clear, far more capable than any we have today) might be able to discover how to create it much more rapidly—and actually help someone do so.

My view is that even though these are obscure risks, and might seem unlikely, the magnitude of the consequences is so large that they should be taken seriously as a first-class risk of AI systems.

Skeptics have raised a number of objections to the seriousness of these biological risks from LLMs, which I disagree with but which are worth addressing. Most fall into the category of not appreciating the exponential trajectory that the technology is on. Back in 2023 when we first started talking about biological risks from LLMs, skeptics said that all the necessary information was available on Google and LLMs didn’t add anything beyond this. It was never true that Google could give you all the necessary information: genomes are freely available, but as I said above, certain key steps, as well as a huge amount of practical know-how cannot be gotten in that way. But also, by the end of 2023 LLMs were clearly providing information beyond what Google could give for some steps of the process.

After this, skeptics retreated to the objection that LLMs weren’t *end-to-end* useful, and couldn’t help with bioweapons *acquisition* as opposed to

Contents

1. I'm sorry, Dave
2. A surprising and terrible empowerment
3. The odious apparatus
4. Player piano
5. Black seas of infinity
- Humanity's test

just providing theoretical information. As of mid-2025, our measurements show that LLMs may already be providing substantial uplift in several relevant areas, perhaps doubling or tripling the likelihood of success. This led to us deciding that Claude Opus 4 (and the subsequent Sonnet 4.5, Opus 4.1, and Opus 4.5 models) needed to be released under our AI Safety Level 3 protections in our Responsible Scaling Policy framework, and to implementing safeguards against this risk (more on this later). We believe that models are likely now approaching the point where, without safeguards, they could be useful in enabling someone with a STEM degree but not specifically a biology degree to go through the whole process of producing a bioweapon.

Another objection is that there are other actions unrelated to AI that society can take to block the production of bioweapons. Most prominently, the gene synthesis industry makes biological specimens on demand, and there is no federal requirement that providers screen orders to make sure they do not contain pathogens. An MIT study found that 36 out of 38 providers fulfilled an order containing the sequence of the 1918 flu. I am supportive of mandated gene synthesis screening that would make it harder for individuals to weaponize pathogens, in order to reduce both AI-driven biological risks and also biological risks in general. But this is not something we have today. It would also be only one tool in reducing risk; it is a complement to guardrails on AI systems, not a substitute.

The best objection is one that I've rarely seen raised: that there is a gap between the models being useful in principle and the actual propensity of bad actors to use them. Most individual bad actors are disturbed individuals, so almost by definition their behavior is unpredictable and irrational—and it's *these* bad actors, the unskilled ones, who might have stood to benefit the most from AI making it much easier to kill many people.²⁴ Just because a type of violent attack is possible, doesn't mean someone will decide to do it. Perhaps biological attacks will be unappealing because they are reasonably likely to infect the perpetrator, they don't cater to the military-style fantasies that many violent individuals or groups have, and it is hard to selectively target specific people. It could also be that going through a process that takes months, even if an AI walks you through it, involves an amount of patience that most disturbed individuals simply don't have. We may simply get lucky and motive and ability don't combine, in practice, in quite the right way.

But this seems like very flimsy protection to rely on. The motives of disturbed loners can change for any reason or no reason, and in fact

Contents

1. I'm sorry, Dave
2. A surprising and terrible empowerment
3. The odious apparatus
4. Player piano
5. Black seas of infinity
- Humanity's test

there are already instances of LLMs being used in attacks (just not with biology). The focus on disturbed loners also ignores ideologically motivated terrorists, who are often willing to expend large amounts of time and effort (for example, the 9/11 hijackers). Wanting to kill as many people as possible is a motive that will probably arise sooner or later, and it unfortunately suggests bioweapons as the method. Even if this motive is extremely rare, it only has to materialize once. And as biology advances (increasingly driven by AI itself), it may also become possible to carry out more selective attacks (for example, targeted against people with specific ancestries), which adds yet another, very chilling, possible motive.

I do not think biological attacks will necessarily be carried out the instant it becomes widely possible to do so—in fact, I would bet against that. But added up across millions of people and a few years of time, I think there is a serious risk of a major attack, and the consequences would be so severe (with casualties potentially in the millions or more) that I believe we have no choice but to take serious measures to prevent it.

Defenses

That brings us to how to defend against these risks. Here I see three things we can do. First, AI companies can put guardrails on their models to prevent them from helping to produce bioweapons. Anthropic is very actively doing this. Claude's Constitution, which mostly focuses on high-level principles and values, has a small number of specific hard-line prohibitions, and one of them relates to helping with the production of biological (or chemical, or nuclear, or radiological) weapons. But all models can be jailbroken, and so as a second line of defense, we've implemented (since mid-2025, when our tests showed our models were starting to get close to the threshold where they might begin to pose a risk) a classifier that specifically detects and blocks bioweapon-related outputs. We regularly upgrade and improve these classifiers, and have generally found them highly robust even against sophisticated adversarial attacks.²⁵ These classifiers increase the costs to serve our models measurably (in some models, they are close to 5% of total inference costs) and thus cut into our margins, but we feel that using them is the right thing to do.

To their credit, some other AI companies have implemented classifiers as well. But not every company has, and there is also nothing requiring companies to keep their classifiers. I am concerned that over time there may be a prisoner's dilemma where companies can defect and lower their costs by removing classifiers. This is once again a classic negative

externalities problem that can't be solved by the voluntary actions of Anthropic or any other single company alone.²⁶ Voluntary industry standards may help, as may third-party evaluations and verification of the type done by AI security institutes and third-party evaluators.

Contents

1. I'm sorry, Dave
2. A surprising and terrible empowerment
3. The odious apparatus
4. Player piano
5. Black seas of infinity
- Humanity's test

But ultimately defense may require government action, which is the second thing we can do. My views here are the same as they are for addressing autonomy risks: we should start with transparency requirements,²⁷ which help society measure, monitor, and collectively defend against risks without disrupting economic activity in a heavy-handed way. Then, if and when we reach clearer thresholds of risk, we can craft legislation that more precisely targets these risks and has a lower chance of collateral damage. In the particular case of bioweapons, I actually think that the time for such targeted legislation may be approaching soon—Anthropic and other companies are learning more and more about the nature of biological risks and what is reasonable to require of companies in defending against them. Fully defending against these risks may require working internationally, even with geopolitical adversaries, but there is precedent in treaties prohibiting the development of biological weapons. I am generally a skeptic about most kinds of international cooperation on AI, but this may be one narrow area where there is some chance of achieving global restraint. Even dictatorships do not want massive bioterrorist attacks.

Finally, the third countermeasure we can take is to try to develop defenses against biological attacks themselves. This could include monitoring and tracking for early detection, investments in air purification R&D (such as far-UVC disinfection), rapid vaccine development that can respond and adapt to an attack, better personal protective equipment (PPE),²⁸ and treatments or vaccinations for some of the most likely biological agents. mRNA vaccines, which can be designed to respond to a particular virus or variant, are an early example of what is possible here. Anthropic is excited to work with biotech and pharmaceutical companies on this problem. But unfortunately I think our expectations on the defensive side should be limited. There is an asymmetry between attack and defense in biology, because agents spread rapidly on their own, while defenses require detection, vaccination, and treatment to be organized across large numbers of people very quickly in response. Unless the response is lightning quick (which it rarely is), much of the damage will be done before a response is possible. It is conceivable that future technological improvements could shift this balance in favor of defense (and we should certainly use AI to help develop such technological advances), but until then, preventative safeguards will be our main line of defense.

Contents

1. I'm sorry, Dave
2. A surprising and terrible empowerment
3. The odious apparatus
4. Player piano
5. Black seas of infinity
- Humanity's test

It's worth a brief mention of cyberattacks here, since unlike biological attacks, AI-led cyberattacks have actually happened in the wild, including at a large scale and for state-sponsored espionage. We expect these attacks to become more capable as models advance rapidly, until they are the main way in which cyberattacks are conducted. I expect AI-led cyberattacks to become a serious and unprecedented threat to the integrity of computer systems around the world, and Anthropic is working very hard to shut down these attacks and eventually reliably prevent them from happening. The reason I haven't focused on cyber as much as biology is that (1) cyberattacks are much less likely to kill people, certainly not at the scale of biological attacks, and (2) the offense-defense balance may be more tractable in cyber, where there is at least some hope that defense could keep up with (and even ideally outpace) AI attack if we invest in it properly.

Although biology is currently the most serious vector of attack, there are many other vectors and it is possible that a more dangerous one may emerge. The general principle is that without countermeasures, AI is likely to continuously lower the barrier to destructive activity on a larger and larger scale, and humanity needs a serious response to this threat.

3. The odious apparatus

Misuse for seizing power

The previous section discussed the risk of individuals and small organizations co-opting a small subset of the “country of geniuses in a datacenter” to cause large-scale destruction. But we should also worry—likely substantially more so—about misuse of AI for the purpose of *wielding or seizing power*, likely by larger and more established actors.²⁹

In *Machines of Loving Grace*, I discussed the possibility that authoritarian governments might use powerful AI to surveil or repress their citizens in ways that would be extremely difficult to reform or overthrow. Current autocracies are limited in how repressive they can be by the need to have humans carry out their orders, and humans often have limits in how inhumane they are willing to be. But AI-enabled autocracies would not have such limits.

Worse yet, countries could also use their advantage in AI to gain power over *other countries*. If the “country of geniuses” as a whole was simply owned and controlled by a single (human) country's military apparatus, and other countries did not have equivalent capabilities, it is hard to see

Contents

1. I'm sorry, Dave
2. A surprising and terrible empowerment
3. The odious apparatus
4. Player piano
5. Black seas of infinity
- Humanity's test

how they could defend themselves: they would be outsmarted at every turn, similar to a war between humans and mice. Putting these two concerns together leads to the alarming possibility of a global totalitarian dictatorship. Obviously, it should be one of our highest priorities to prevent this outcome.

There are many ways in which AI could enable, entrench, or expand autocracy, but I'll list a few that I'm most worried about. Note that some of these applications have legitimate defensive uses, and I am not necessarily arguing against them in absolute terms; I am nevertheless worried that they structurally tend to favor autocracies:

- **Fully autonomous weapons.** A swarm of millions or billions of fully automated armed drones, locally controlled by powerful AI and strategically coordinated across the world by an even more powerful AI, could be an unbeatable army, capable of both defeating any military in the world and suppressing dissent within a country by following around every citizen. Developments in the Russia-Ukraine War should alert us to the fact that drone warfare is already with us (though not fully autonomous yet, and a tiny fraction of what might be possible with powerful AI). R&D from powerful AI could make the drones of one country far superior to those of others, speed up their manufacture, make them more resistant to electronic attacks, improve their maneuvering, and so on. Of course, these weapons also have legitimate uses in the defense of democracy: they have been key to defending Ukraine and would likely be key to defending Taiwan. But they are a dangerous weapon to wield: we should worry about them in the hands of autocracies, but also worry that because they are so powerful, with so little accountability, there is a greatly increased risk of democratic governments turning them against their own people to seize power.
- **AI surveillance.** Sufficiently powerful AI could likely be used to compromise any computer system in the world,³⁰ and could also use the access obtained in this way to read *and make sense of* all the world's electronic communications (or even all the world's in-person communications, if recording devices can be built or commandeered). It might be frighteningly plausible to simply generate a complete list of anyone who disagrees with the government on any number of issues, even if such disagreement isn't explicit in anything they say or do. A powerful AI looking across billions of conversations from millions of people could gauge public sentiment, detect pockets of disloyalty forming, and stamp them out before they grow. This could lead to the

Contents

1. I'm sorry, Dave
2. A surprising and terrible empowerment
3. The odious apparatus
4. Player piano
5. Black seas of infinity
- Humanity's test

imposition of a true panopticon on a scale that we don't see today, even with the CCP.

- **AI propaganda.** Today's phenomena of "AI psychosis" and "AI girlfriends" suggest that even at their current level of intelligence, AI models can have a powerful psychological influence on people. Much more powerful versions of these models, that were much more embedded in and aware of people's daily lives and could model and influence them over months or years, would likely be capable of essentially brainwashing many (most?) people into any desired ideology or attitude, and could be employed by an unscrupulous leader to ensure loyalty and suppress dissent, even in the face of a level of repression that most populations would rebel against. Today people worry a lot about, for example, the potential influence of TikTok as CCP propaganda directed at children. I worry about that too, but a personalized AI agent that gets to know you over years and uses its knowledge of you to shape all of your opinions would be dramatically more powerful than this.
- **Strategic decision-making.** A country of geniuses in a datacenter could be used to advise a country, group, or individual on geopolitical strategy, what we might call a "virtual Bismarck." It could optimize the three strategies above for seizing power, plus probably develop many others that I haven't thought of (but that a country of geniuses could). Diplomacy, military strategy, R&D, economic strategy, and many other areas are all likely to be substantially increased in effectiveness by powerful AI. Many of these skills would be legitimately helpful for democracies—we want democracies to have access to the best strategies for defending themselves against autocracies—but the potential for misuse in *anyone's* hands still remains.

Having described *what* I am worried about, let's move on to *who*. I am worried about entities who have the most access to AI, who are starting from a position of the most political power, or who have an existing history of repression. In order of severity, I am worried about:

- **The CCP.** China is second only to the United States in AI capabilities, and is the country with the greatest likelihood of surpassing the United States in those capabilities. Their government is currently autocratic and operates a high-tech surveillance state. It has deployed AI-based surveillance already (including in the repression of Uyghurs), and is believed to employ algorithmic propaganda via TikTok (in addition to its many other international propaganda efforts). They have hands down the

Contents

1. I'm sorry, Dave
 2. A surprising and terrible empowerment
 3. The odious apparatus
 4. Player piano
 5. Black seas of infinity
- Humanity's test

clearest path to the AI-enabled totalitarian nightmare I laid out above. It may even be the default outcome within China, as well as within other autocratic states to whom the CCP exports surveillance technology. I have written often about the threat of the CCP taking the lead in AI and the existential imperative to prevent them from doing so. This is why. To be clear, I am not singling out China out of animus to them in particular—they are simply the country that most combines AI prowess, an autocratic government, and a high-tech surveillance state. If anything, it is the Chinese people themselves who are most likely to suffer from the CCP's AI-enabled repression, and they have no voice in the actions of their government. I greatly admire and respect the Chinese people and support the many brave dissidents within China and their struggle for freedom.

- **Democracies competitive in AI.** As I wrote above, democracies have a legitimate interest in some AI-powered military and geopolitical tools, because democratic governments offer the best chance to counter the use of these tools by autocracies. Broadly, I am supportive of arming democracies with the tools needed to defeat autocracies in the age of AI—I simply don't think there is any other way. But we cannot ignore the potential for abuse of these technologies by democratic governments themselves. Democracies normally have safeguards that prevent their military and intelligence apparatus from being turned inwards against their own population,³¹ but because AI tools require so few people to operate, there is potential for them to circumvent these safeguards and the norms that support them. It is also worth noting that some of these safeguards are already gradually eroding in some democracies. Thus, we should arm democracies with AI, but we should do so carefully and within limits: they are the immune system we need to fight autocracies, but like the immune system, there is some risk of them turning on us and becoming a threat themselves.
- **Non-democratic countries with large datacenters.** Beyond China, most countries with less democratic governance are not leading AI players in the sense that they don't have companies which produce frontier AI models. Thus they pose a fundamentally different and lesser risk than the CCP, which remains the primary concern (most are also less repressive, and the ones that are more repressive, like North Korea, have no significant AI industry at all). But some of these countries do have large *datacenters* (often as part of buildouts by companies operating in democracies), which can be used to run frontier AI at

Contents

1. I'm sorry, Dave
2. A surprising and terrible empowerment
3. The odious apparatus
4. Player piano
5. Black seas of infinity
- Humanity's test

large scale (though this does not confer the ability to push the frontier). There is some amount of danger associated with this—these governments could in principle expropriate the datacenters and use the country of AIs within it for their own ends. I am less worried about this compared to countries like China that directly develop AI, but it's a risk to keep in mind. ³²

- **AI companies.** It is somewhat awkward to say this as the CEO of an AI company, but I think the next tier of risk is actually AI companies themselves. AI companies control large datacenters, train frontier models, have the greatest expertise on how to use those models, and in some cases have daily contact with and the possibility of influence over tens or hundreds of millions of users. The main thing they lack is the legitimacy and infrastructure of a state, so much of what would be needed to build the tools of an AI autocracy would be illegal for an AI company to do, or at least exceedingly suspicious. But some of it is not impossible: they could, for example, use their AI products to brainwash their massive consumer user base, and the public should be alert to the risk this represents. I think the governance of AI companies deserves a lot of scrutiny.

There are a number of possible arguments against the severity of these threats, and I wish I believed them, because AI-enabled authoritarianism terrifies me. It's worth going through some of these arguments and responding to them.

First, some people might put their faith in the nuclear deterrent, particularly to counter the use of AI autonomous weapons for military conquest. If someone threatens to use these weapons against you, you can always threaten a nuclear response back. My worry is that I'm not totally sure we can be confident in the nuclear deterrent against a country of geniuses in a datacenter: it is possible that powerful AI could devise ways to detect and strike nuclear submarines, conduct influence operations against the operators of nuclear weapons infrastructure, or use AI's cyber capabilities to launch a cyberattack against satellites used to detect nuclear launches. ³³ Alternatively, it's possible that taking over countries is feasible with only AI surveillance and AI propaganda, and never actually presents a clear moment where it's obvious what is going on and where a nuclear response would be appropriate. *Maybe* these things aren't feasible and the nuclear deterrent will still be effective, but it seems too high stakes to take a risk. ³⁴

Contents

1. I'm sorry, Dave
2. A surprising and terrible empowerment
3. The odious apparatus
4. Player piano
5. Black seas of infinity
- Humanity's test

A second possible objection is that there might be countermeasures we can take against these tools of autocracy. We can counter drones with our own drones, cyberdefense will improve along with cyberattack, there may be ways to immunize people against propaganda, etc. My response is that these defenses will only be possible with comparably powerful AI. If there isn't some counterforce with a comparably smart and numerous country of geniuses in a datacenter, it won't be possible to match the quality or quantity of drones, for cyberdefense to outsmart cyberoffense, etc. So the question of countermeasures reduces to the question of a balance of power in powerful AI. Here, I am concerned about the recursive or self-reinforcing property of powerful AI (which I discussed at the beginning of this essay): that each generation of AI can be used to design and train the next generation of AI. This leads to a risk of a runaway advantage, where the current leader in powerful AI may be able to increase their lead and may be difficult to catch up with. We need to make sure it is not an authoritarian country that gets to this loop first.

Furthermore, even if a balance of power can be achieved, there is still risk that the world could be split up into autocratic spheres, as in *Nineteen Eighty-Four*. Even if several competing powers each have their powerful AI models, and none can overpower the others, each power could still internally repress their own population, and would be very difficult to overthrow (since the populations don't have powerful AI to defend themselves). It is thus important to prevent AI-enabled autocracy even if it doesn't lead to a single country taking over the world.

Defenses

How do we defend against this wide range of autocratic tools and potential threat actors? As in the previous sections, there are several things I think we can do. First, we should absolutely not be selling chips, chip-making tools, or datacenters to the CCP. Chips and chip-making tools are the single greatest bottleneck to powerful AI, and blocking them is a simple but extremely effective measure, perhaps the most important single action we can take. It makes no sense to sell the CCP the tools with which to build an AI totalitarian state and possibly conquer us militarily. A number of complicated arguments are made to justify such sales, such as the idea that "spreading our tech stack around the world" allows "America to win" in some general, unspecified economic battle. In my view, this is like selling nuclear weapons to North Korea and then bragging that the missile casings are made by Boeing and so the US is "winning." China is several years behind the

US in their ability to produce frontier chips in quantity, and the critical period for building the country of geniuses in a datacenter is very likely to be within those next several years.³⁵ There is no reason to give a giant boost to their AI industry during this critical period.

Contents

1. I'm sorry, Dave
2. A surprising and terrible empowerment
3. The odious apparatus
4. Player piano
5. Black seas of infinity
- Humanity's test

Second, it makes sense to use AI to empower democracies to resist autocracies. This is the reason Anthropic considers it important to provide AI to the intelligence and defense communities in the US and its democratic allies. Defending democracies that are under attack, such as Ukraine and (via cyber attacks) Taiwan, seems especially high priority, as does empowering democracies to use their intelligence services to disrupt and degrade autocracies from the inside. At some level the only way to respond to autocratic threats is to match and outclass them militarily. A coalition of the US and its democratic allies, if it achieved predominance in powerful AI, would be in a position to not only defend itself against autocracies, but contain them and limit their AI totalitarian abuses.

Third, we need to draw a hard line against AI abuses within democracies. There need to be limits to what we allow our governments to do with AI, so that they don't seize power or repress their own people. The formulation I have come up with is that we should use AI for national defense in all ways *except those which would make us more like our autocratic adversaries*.

Where should the line be drawn? In the list at the beginning of this section, two items—using AI for domestic mass surveillance and mass propaganda—seem to me like bright red lines and entirely illegitimate. Some might argue that there's no need to do anything (at least in the US), since domestic mass surveillance is already illegal under the Fourth Amendment. But the rapid progress of AI may create situations that our existing legal frameworks are not well designed to deal with. For example, it would likely not be unconstitutional for the US government to conduct massively scaled recordings of all *public* conversations (e.g., things people say to each other on a street corner), and previously it would have been difficult to sort through this volume of information, but with AI it could all be transcribed, interpreted, and triangulated to create a picture of the attitude and loyalties of many or most citizens. I would support civil liberties-focused legislation (or maybe even a constitutional amendment) that imposes stronger guardrails against AI-powered abuses.

The other two items—fully autonomous weapons and AI for strategic decision-making—are harder lines to draw since they have legitimate uses in defending democracy, while also being prone to abuse. Here I

Contents

1. I'm sorry, Dave
2. A surprising and terrible empowerment
3. The odious apparatus
4. Player piano
5. Black seas of infinity
- Humanity's test

think what is warranted is extreme care and scrutiny combined with guardrails to prevent abuses. My main fear is having too small a number of “fingers on the button,” such that one or a handful of people could essentially operate a drone army without needing any other humans to cooperate to carry out their orders. As AI systems get more powerful, we may need to have more direct and immediate oversight mechanisms to ensure they are not misused, perhaps involving branches of government other than the executive. I think we should approach fully autonomous weapons in particular with great caution,³⁶ and not rush into their use without proper safeguards.

Fourth, after drawing a hard line against AI abuses in democracies, we should use that precedent to create an international taboo against the worst abuses of powerful AI. I recognize that the current political winds have turned against international cooperation and international norms, but this is a case where we sorely need them. The world needs to understand the dark potential of powerful AI in the hands of autocrats, and to recognize that certain uses of AI amount to an attempt to permanently steal their freedom and impose a totalitarian state from which they can't escape. I would even argue that in some cases, large-scale surveillance with powerful AI, mass propaganda with powerful AI, and certain types of *offensive* uses of fully autonomous weapons should be considered crimes against humanity. More generally, a robust norm against AI-enabled totalitarianism and all its tools and instruments is sorely needed.

It is possible to have an even stronger version of this position, which is that because the possibilities of AI-enabled totalitarianism are so dark, autocracy is simply not a form of government that people can accept in the post-powerful AI age. Just as feudalism became unworkable with the industrial revolution, the AI age could lead inevitably and logically to the conclusion that democracy (and, hopefully, democracy improved and reinvigorated by AI, as I discuss in *Machines of Loving Grace*) is the only viable form of government if humanity is to have a good future.

Fifth and finally, AI companies should be carefully watched, as should their connection to the government, which is necessary, but must have limits and boundaries. The sheer amount of capability embodied in powerful AI is such that ordinary corporate governance—which is designed to protect shareholders and prevent ordinary abuses such as fraud—is unlikely to be up to the task of governing AI companies. There may also be value in companies publicly committing to (perhaps even as part of corporate governance) not take certain actions, such as privately building or stockpiling military hardware, using large

amounts of computing resources by single individuals in unaccountable ways, or using their AI products as propaganda to manipulate public opinion in their favor.

The danger here comes from many directions, and some directions are in tension with others. The only constant is that we must seek accountability, norms, and guardrails for everyone, even as we empower “good” actors to keep “bad” actors in check.

Contents

1. I’m sorry, Dave
2. A surprising and terrible empowerment
3. The odious apparatus
4. Player piano
5. Black seas of infinity
- Humanity’s test

4. Player piano

Economic disruption

The previous three sections were essentially about security risks posed by powerful AI: risks from the AI itself, risks from misuse by individuals and small organizations and risks of misuse by states and large organizations. If we put aside security risks or assume they have been solved, the next question is economic. What will be the effect of this infusion of incredible “human” capital on the economy? Clearly, the most obvious effect will be to greatly increase economic growth. The pace of advances in scientific research, biomedical innovation, manufacturing, supply chains, the efficiency of the financial system, and much more are almost guaranteed to lead to a much faster rate of economic growth. In *Machines of Loving Grace*, I suggest that a 10–20% sustained annual GDP growth rate may be possible.

But it should be clear that this is a double-edged sword: what are the economic prospects for most existing humans in such a world? New technologies often bring labor market shocks, and in the past humans have always recovered from them, but I am concerned that this is because these previous shocks affected only a small fraction of the full possible range of human abilities, leaving room for humans to expand to new tasks. AI will have effects that are much broader and occur much faster, and therefore I worry it will be much more challenging to make things work out well.

Labor market disruption

There are two specific problems I am worried about: labor market displacement, and concentration of economic power. Let’s start with the first one. This is a topic that I warned about very publicly in 2025, where I predicted that AI could displace half of all entry-level white collar jobs in the next 1–5 years, even as it accelerates economic growth and scientific progress. This warning started a public debate about the topic. Many CEOs, technologists, and economists agreed with me, but

Contents

1. I'm sorry, Dave
2. A surprising and terrible empowerment
3. The odious apparatus
4. Player piano
5. Black seas of infinity
- Humanity's test

others assumed I was falling prey to a “lump of labor” fallacy and didn’t know how labor markets worked, and some didn’t see the 1–5-year time range and thought I was claiming AI is displacing jobs right now (which I agree it is likely not). So it is worth going through in detail why I am worried about labor displacement, to clear up these misunderstandings.

As a baseline, it’s useful to understand how labor markets *normally* respond to advances in technology. When a new technology comes along, it starts by making pieces of a given human job more efficient. For example, early in the Industrial Revolution, machines, such as upgraded plows, enabled human farmers to be more efficient at some aspects of the job. This improved the productivity of farmers, which increased their wages.

In the next step, some parts of the job of farming could be done *entirely* by machines, for example with the invention of the threshing machine or seed drill. In this phase, humans did a lower and lower fraction of the job, but the work they *did* complete became more and more leveraged because it is complementary to the work of machines, and their productivity continued to rise. As described by Jevons’ paradox, the wages of farmers and perhaps even the number of farmers continued to increase. Even when 90% of the job is being done by machines, humans can simply do 10x more of the 10% they still do, producing 10x as much output for the same amount of labor.

Eventually, machines do everything or almost everything, as with modern combine harvesters, tractors, and other equipment. At this point farming as a form of human employment really does go into steep decline, and this potentially causes serious disruption in the short term, but because farming is just one of many useful activities that humans are able to do, people eventually switch to other jobs, such as operating factory machines. This is true even though farming accounted for a huge proportion of employment *ex ante*. 250 years ago, 90% of Americans lived on farms; in Europe, 50–60% of employment was agricultural. Now those percentages are in the low single digits in those places, because workers switched to industrial jobs (and later, knowledge work jobs). The economy can do what previously required most of the labor force with only 1–2% of it, freeing up the rest of the labor force to build an ever more advanced industrial society. There’s no fixed “lump of labor,” just an ever-expanding ability to do more and more with less and less. People’s wages rise in line with the GDP exponential and the economy maintains full employment once disruptions in the short term have passed.

It's possible things will go roughly the same way with AI, but I would bet pretty strongly against it. Here are some reasons I think AI is likely to be different:

Contents

1. I'm sorry, Dave
2. A surprising and terrible empowerment
3. The odious apparatus
4. Player piano
5. Black seas of infinity
- Humanity's test

- **Speed.** The pace of progress in AI is much faster than for previous technological revolutions. For example, in the last 2 years, AI models went from barely being able to complete a single line of code, to writing all or almost all of the code for some people—including engineers at Anthropic.³⁷ Soon, they may do the entire task of a software engineer end to end.³⁸ It is hard for people to adapt to this pace of change, both to the changes in how a given job works and in the need to switch to new jobs. Even legendary programmers are increasingly describing themselves as “behind.” The pace may if anything continue to speed up, as AI coding models increasingly accelerate the task of AI development. To be clear, speed in itself does not mean labor markets and employment won't eventually recover, it just implies the short-term transition will be unusually painful compared to past technologies, since humans and labor markets are slow to react and to equilibrate.
- **Cognitive breadth.** As suggested by the phrase “country of geniuses in a datacenter,” AI will be capable of a very wide range of human cognitive abilities—perhaps all of them. This is very different from previous technologies like mechanized farming, transportation, or even computers.³⁹ This will make it harder for people to switch easily from jobs that are displaced to similar jobs that they would be a good fit for. For example, the general intellectual abilities required for entry-level jobs in, say, finance, consulting, and law are fairly similar, even if the specific knowledge is quite different. A technology that disrupted only one of these three would allow employees to switch to the two other close substitutes (or for undergraduates to switch majors). But disrupting all three at once (along with many other similar jobs) may be harder for people to adapt to. Furthermore, it's not *just* that most existing jobs will be disrupted. That part has happened before—recall that farming was a huge percentage of employment. But farmers could switch to the relatively similar work of operating factory machines, even though that work hadn't been common before. By contrast, AI is increasingly matching the general cognitive profile of humans, which means it will also be good at the new jobs that would ordinarily be created in response to the old ones being automated. Another way to say it is that AI

Contents

1. I'm sorry, Dave
 2. A surprising and terrible empowerment
 3. The odious apparatus
 4. Player piano
 5. Black seas of infinity
- Humanity's test

isn't a substitute for specific human jobs but rather a general labor substitute for humans.

- **Slicing by cognitive ability.** Across a wide range of tasks, AI appears to be advancing from the bottom of the ability ladder to the top. For example, in coding our models have proceeded from the level of “a mediocre coder” to “a strong coder” to “a very strong coder.”⁴⁰ We are now starting to see the same progression in white-collar work in general. We are thus at risk of a situation where, instead of affecting people with specific skills or in specific professions (who can adapt by retraining), AI is affecting people with certain intrinsic cognitive properties, namely lower intellectual ability (which is harder to change). It is not clear where these people will go or what they will do, and I am concerned that they could form an unemployed or very-low-wage “underclass.” To be clear, things somewhat like this have happened before—for example, computers and the internet are believed by some economists to represent “skill-biased technological change.” But this skill biasing was both not as extreme as what I expect to see with AI, and is believed to have contributed to an increase in wage inequality,⁴¹ so it is not exactly a reassuring precedent.
- **Ability to fill in the gaps.** The way human jobs often adjust in the face of new technology is that there are many aspects to the job, and the new technology, even if it appears to directly replace humans, often has gaps in it. If someone invents a machine to make widgets, humans may still have to load raw material into the machine. Even if that takes only 1% as much effort as making the widgets manually, human workers can simply make 100x more widgets. But AI, in addition to being a rapidly advancing technology, is also a rapidly *adapting* technology. During every model release, AI companies carefully measure what the model is good at and what it isn't, and customers also provide such information after the launch. Weaknesses can be addressed by collecting tasks that embody the current gap, and training on them for the next model. Early in generative AI, users noticed that AI systems had certain weaknesses (such as AI image models generating hands with the wrong number of fingers) and many assumed these weaknesses were inherent to the technology. If they were, it would limit job disruption. But pretty much every such weakness gets addressed quickly— often, within just a few months.

It's worth addressing common points of skepticism. First, there is the argument that economic diffusion will be slow, such that even if the underlying technology is *capable* of doing most human labor, the actual

Contents

1. I'm sorry, Dave
2. A surprising and terrible empowerment
3. The odious apparatus
4. Player piano
5. Black seas of infinity
- Humanity's test

application of it across the economy may be much slower (for example in industries that are far from the AI industry and slow to adopt). Slow diffusion of technology is definitely real—I talk to people from a wide variety of enterprises, and there are places where the adoption of AI will take years. That's why my prediction for 50% of entry level white collar jobs being disrupted is 1–5 years, even though I suspect we'll have powerful AI (which would be, technologically speaking, enough to do *most or all* jobs, not just entry level) in much less than 5 years. But diffusion effects merely buy us time. And I am not confident they will be as slow as people predict. Enterprise AI adoption is growing at rates much faster than any previous technology, largely on the pure strength of the technology itself. Also, even if traditional enterprises are slow to adopt new technology, startups will spring up to serve as “glue” and make the adoption easier. If that doesn't work, the startups may simply disrupt the incumbents directly.

That could lead to a world where it isn't so much that specific jobs are disrupted as it is that large enterprises are disrupted in general and replaced with much less labor-intensive startups. This could also lead to a world of “geographic inequality,” where an increasing fraction of the world's wealth is concentrated in Silicon Valley, which becomes its own economy running at a different speed than the rest of the world and leaving it behind. All of these outcomes would be great for economic growth—but not so great for the labor market or those who are left behind.

Second, some people say that human jobs will move to the physical world, which avoids the whole category of “cognitive labor” where AI is progressing so rapidly. I am not sure how safe this is, either. A lot of physical labor is already being done by machines (e.g., manufacturing) or will soon be done by machines (e.g., driving). Also, sufficiently powerful AI will be able to accelerate the development of robots, and then control those robots in the physical world. It may buy some time (which is a good thing), but I'm worried it won't buy much. And even if the disruption was limited only to cognitive tasks, it would still be an unprecedentedly large and rapid disruption.

Third, perhaps some tasks inherently require or greatly benefit from a human touch. I'm a little more uncertain about this one, but I'm still skeptical that it will be enough to offset the bulk of the impacts I described above. AI is already widely used for customer service. Many people report that it is easier to talk to AI about their personal problems than to talk to a therapist—that the AI is more patient. When my sister was struggling with medical problems during a pregnancy,

Contents

1. I'm sorry, Dave
2. A surprising and terrible empowerment
3. The odious apparatus
4. Player piano
5. Black seas of infinity
- Humanity's test

she felt she wasn't getting the answers or support she needed from her care providers, and she found Claude to have a better bedside manner (as well as succeeding better at diagnosing the problem). I'm sure there are some tasks for which a human touch really is important, but I'm not sure how many—and here we're talking about finding work for nearly everyone in the labor market.

Fourth, some may argue that comparative advantage will still protect humans. Under the law of comparative advantage, even if AI is better than humans at everything, any *relative* differences between the human and AI profile of skills creates a basis of trade and specialization between humans and AI. The problem is that if AIs are literally thousands of times more productive than humans, this logic starts to break down. Even tiny transaction costs could make it not worth it for AI to trade with humans. And human wages may be very low, even if they technically have something to offer.

It's possible all of these factors can be addressed—that the labor market is resilient enough to adapt to even such an enormous disruption. But even if it can eventually adapt, the factors above suggest that the short-term shock will be unprecedented in size.

Defenses

What can we do about this problem? I have several suggestions, some of which Anthropic is already doing. The first thing is simply to get accurate data about what is happening with job displacement in real time. When an economic change happens very quickly, it's hard to get reliable data about what is happening, and without reliable data it is hard to design effective policies. For example, government data is currently lacking granular, high-frequency data on AI adoption across firms and industries. For the last year Anthropic has been operating and publicly releasing an Economic Index that shows use of our models almost in real time, broken down by industry, task, location, and even things like whether a task was being automated or conducted collaboratively. We also have an Economic Advisory Council to help us interpret this data and see what is coming.

Second, AI companies have a choice in how they work with enterprises. The very inefficiency of traditional enterprises means that their rollout of AI can be very path dependent, and there is some room to choose a better path. Enterprises often have a choice between “cost savings” (doing the same thing with fewer people) and “innovation” (doing more with the same number of people). The market will inevitably produce both eventually, and any competitive AI company will have to serve

some of both, but there may be some room to steer companies towards innovation when possible, and it may buy us some time. Anthropic is actively thinking about this.

Contents

1. I'm sorry, Dave
2. A surprising and terrible empowerment
3. The odious apparatus
4. Player piano
5. Black seas of infinity
- Humanity's test

Third, companies should think about how to take care of their employees. In the short term, being creative about ways to reassign employees within companies may be a promising way to stave off the need for layoffs. In the long term, in a world with enormous total wealth, in which many companies increase greatly in value due to increased productivity and capital concentration, it may be feasible to pay human employees even long after they are no longer providing economic value in the traditional sense. Anthropic is currently considering a range of possible pathways for our own employees that we will share in the near future.

Fourth, wealthy individuals have an obligation to help solve this problem. It is sad to me that many wealthy individuals (especially in the tech industry) have recently adopted a cynical and nihilistic attitude that philanthropy is inevitably fraudulent or useless. Both private philanthropy like the [Gates Foundation](#) and public programs like [PEPFAR](#) have saved tens of millions of lives in the developing world, and helped to create economic opportunity in the developed world. All of Anthropic's co-founders have pledged to donate 80% of our wealth, and Anthropic's staff have individually pledged to donate company shares worth billions at current prices—donations that the company has committed to matching.

Fifth, while all the above private actions can be helpful, ultimately a macroeconomic problem this large will require government intervention. The natural policy response to an enormous economic pie coupled with high inequality (due to a lack of jobs, or poorly paid jobs, for many) is progressive taxation. The tax could be general or could be targeted against AI companies in particular. Obviously tax design is complicated, and there are many ways for it to go wrong. I don't support poorly designed tax policies. I think the extreme levels of inequality predicted in this essay justify a more robust tax policy on basic moral grounds, but I can also make a pragmatic argument to the world's billionaires that it's in their interest to support a good version of it: if they don't support a good version, they'll inevitably get a bad version designed by a mob.

Ultimately, I think of all of the above interventions as ways to buy time. In the end AI will be able to do everything, and we need to grapple with that. It's my hope that by that time, we can use AI itself to help us

restructure markets in ways that work for everyone, and that the interventions above can get us through the transitional period.

Economic concentration of power

Contents

1. I'm sorry, Dave
2. A surprising and terrible empowerment
3. The odious apparatus
4. Player piano
5. Black seas of infinity
- Humanity's test

Separate from the problem of job displacement or economic inequality *per se* is the problem of *economic concentration of power*. Section 1 discussed the risk that humanity gets disempowered by AI, and Section 3 discussed the risk that citizens get disempowered by their governments by force or coercion. But another kind of disempowerment can occur if there is such a huge concentration of wealth that a small group of people effectively controls government policy with their influence, and ordinary citizens have no influence because they lack economic leverage. Democracy is ultimately backstopped by the idea that the population as a whole is necessary for the operation of the economy. If that economic leverage goes away, then the implicit social contract of democracy may stop working. Others have written about this, so I needn't go into great detail about it here, but I agree with the concern, and I worry it is already starting to happen.

To be clear, I am not opposed to people making a lot of money. There's a strong argument that it incentivizes economic growth under normal conditions. I am sympathetic to concerns about impeding innovation by killing the golden goose that generates it. But in a scenario where GDP growth is 10–20% a year and AI is rapidly taking over the economy, yet single individuals hold appreciable fractions of the GDP, innovation is *not* the thing to worry about. The thing to worry about is a level of wealth concentration that will break society.

The most famous example of extreme concentration of wealth in US history is the Gilded Age, and the wealthiest industrialist of the Gilded Age was John D. Rockefeller. Rockefeller's wealth amounted to ~2% of the US GDP at the time.⁴² A similar fraction today would lead to a fortune of \$600B, and the richest person in the world today (Elon Musk) already exceeds that, at roughly \$700B. So we are already at historically unprecedented levels of wealth concentration, even *before* most of the economic impact of AI. I don't think it is too much of a stretch (if we get a "country of geniuses") to imagine AI companies, semiconductor companies, and perhaps downstream application companies generating ~\$3T in revenue per year,⁴³ being valued at ~\$30T, and leading to personal fortunes well into the trillions. In that world, the debates we have about tax policy today simply won't apply as we will be in a fundamentally different situation.

Contents

1. I'm sorry, Dave
2. A surprising and terrible empowerment
3. The odious apparatus
4. Player piano
5. Black seas of infinity
- Humanity's test

Related to this, the coupling of this economic concentration of wealth with the political system already concerns me. AI datacenters already represent a substantial fraction of US economic growth,⁴⁴ and are thus strongly tying together the financial interests of large tech companies (which are increasingly focused on either AI or AI infrastructure) and the political interests of the government in a way that can produce perverse incentives. We already see this through the reluctance of tech companies to criticize the US government, and the government's support for extreme anti-regulatory policies on AI.

Defenses

What can be done about this? First, and most obviously, companies should simply choose not to be part of it. Anthropic has always strived to be a policy actor and not a political one, and to maintain our authentic views whatever the administration. We've spoken up in favor of sensible AI regulation and export controls that are in the public interest, even when these are at odds with government policy.⁴⁵ Many people have told me that we should stop doing this, that it could lead to unfavorable treatment, but in the year we've been doing it, Anthropic's valuation has increased by over 6x, an almost unprecedented jump at our commercial scale.

Second, the AI industry needs a healthier relationship with government—one based on substantive policy engagement rather than political alignment. Our choice to engage on policy substance rather than politics is sometimes read as a tactical error or failure to “read the room” rather than a principled decision, and that framing concerns me. In a healthy democracy, companies should be able to advocate for good policy for its own sake. Related to this, a public backlash against AI is brewing: this could be a corrective, but it's currently unfocused. Much of it targets issues that aren't actually problems (like datacenter water usage) and proposes solutions (like datacenter bans or poorly designed wealth taxes) that wouldn't address the real concerns. The underlying issue that deserves attention is ensuring that AI development remains accountable to the public interest, not captured by any particular political or commercial alliance, and it seems important to focus the public discussion there.

Third, the macroeconomic interventions I described earlier in this section, as well as a resurgence of private philanthropy, can help to balance the economic scales, addressing both the job displacement and concentration of economic power problems at once. We should look to the history of our country here: even in the Gilded Age, industrialists such as Rockefeller and Carnegie felt a strong obligation to society at

Contents

1. I'm sorry, Dave
2. A surprising and terrible empowerment
3. The odious apparatus
4. Player piano
5. Black seas of infinity
- Humanity's test

large, a feeling that society had contributed enormously to their success and they needed to give back. That spirit seems to be increasingly missing today, and I think it is a large part of the way out of this economic dilemma. Those who are at the forefront of AI's economic boom should be willing to give away both their wealth and their power.

5. Black seas of infinity

Indirect effects

This last section is a catchall for unknown unknowns, particularly things that could go wrong as an indirect result of positive advances in AI and the resulting acceleration of science and technology in general. Suppose we address all the risks described so far, and begin to reap the benefits of AI. We will likely get a “century of scientific and economic progress compressed into a decade,” and this will be hugely positive for the world, but we will then have to contend with the problems that arise from this rapid rate of progress, and those problems may come at us fast. We may also encounter other risks that occur indirectly as a consequence of AI progress and are hard to anticipate in advance.

By the nature of unknown unknowns it is impossible to make an exhaustive list, but I'll list three possible concerns as illustrative examples for what we should be watching for:

- **Rapid advances in biology.** If we do get a century of medical progress in a few years, it is possible that we will greatly increase the human lifespan, and there is a chance we also gain radical capabilities like the ability to increase human intelligence or radically modify human biology. Those would be big changes in what is possible, happening very quickly. They could be positive if responsibly done (which is my hope, as described in *Machines of Loving Grace*), but there is always a risk they go very wrong—for example, if efforts to make humans smarter also make them more unstable or power-seeking. There is also the issue of “uploads” or “whole brain emulation,” digital human minds instantiated in software, which might someday help humanity transcend its physical limitations, but which also carry risks I find disquieting.
- **AI changes human life in an unhealthy way.** A world with billions of intelligences that are much smarter than humans at everything is going to be a very weird world to live in. Even if AI doesn't actively aim to attack humans (Section 1), and isn't explicitly used for oppression or control by states (Section 3), there is a lot that could go wrong short of this, via normal business

Contents

1. I'm sorry, Dave
2. A surprising and terrible empowerment
3. The odious apparatus
4. Player piano
5. Black seas of infinity
- Humanity's test

incentives and nominally consensual transactions. We see early hints of this in the concerns about AI psychosis, AI driving people to suicide, and concerns about romantic relationships with AIs. As an example, could powerful AIs invent some new religion and convert millions of people to it? Could most people end up “addicted” in some way to AI interactions? Could people end up being “puppeted” by AI systems, where an AI essentially watches their every move and tells them exactly what to do and say at all times, leading to a “good” life but one that lacks freedom or any pride of accomplishment? It would not be hard to generate dozens of these scenarios if I sat down with the creator of *Black Mirror* and tried to brainstorm them. I think this points to the importance of things like improving Claude's Constitution, over and above what is necessary for preventing the issues in Section 1. Making sure that AI models *really* have their users' long-term interests at heart, in a way thoughtful people would endorse rather than in some subtly distorted way, seems critical.

- **Human purpose.** This is related to the previous point, but it's not so much about specific human interactions with AI systems as it is about how human life changes in general in a world with powerful AI. Will humans be able to find purpose and meaning in such a world? I think this is a matter of attitude: as I said in *Machines of Loving Grace*, I think human purpose does not depend on being the best in the world at something, and humans can find purpose even over very long periods of time through stories and projects that they love. We simply need to break the link between the generation of economic value and self-worth and meaning. But that is a transition society has to make, and there is always the risk we don't handle it well.

My hope with all of these potential problems is that in a world with powerful AI that we trust not to kill us, that is not the tool of an oppressive government, and that is genuinely working on our behalf, we can use AI itself to anticipate and prevent these problems. But that is not guaranteed—like all of the other risks, it is something we have to handle with care.

Humanity's test

Reading this essay may give the impression that we are in a daunting situation. I certainly found it daunting to write, in contrast with *Machines of Loving Grace*, which felt like giving form and structure to surpassingly beautiful music that had been echoing in my head for

Contents

1. I'm sorry, Dave
2. A surprising and terrible empowerment
3. The odious apparatus
4. Player piano
5. Black seas of infinity
- Humanity's test

years. And there is much about the situation that genuinely is hard. AI brings threats to humanity from multiple directions, and there is genuine tension between the different dangers, where mitigating some of them risks making others worse if we do not thread the needle extremely carefully.

Taking time to carefully build AI systems so they do not autonomously threaten humanity is in genuine tension with the need for democratic nations to stay ahead of authoritarian nations and not be subjugated by them. But in turn, the same AI-enabled tools that are necessary to fight autocracies can, if taken too far, be turned inward to create tyranny in our own countries. AI-driven terrorism could kill millions through the misuse of biology, but an overreaction to this risk could lead us down the road to an autocratic surveillance state. The labor and economic concentration effects of AI, in addition to being grave problems in their own right, may force us to face the other problems in an environment of public anger and perhaps even civil unrest, rather than being able to call on the better angels of our nature. Above all, the sheer *number* of risks, including unknown ones, and the need to deal with all of them at once, creates an intimidating gauntlet that humanity must run.

Furthermore, the last few years should make clear that the idea of stopping or even substantially slowing the technology is fundamentally untenable. The formula for building powerful AI systems is incredibly simple, so much so that it can almost be said to emerge spontaneously from the right combination of data and raw computation. Its creation was probably inevitable the instant humanity invented the transistor, or arguably even earlier when we first learned to control fire. If one company does not build it, others will do so nearly as fast. If all companies in democratic countries stopped or slowed development, by mutual agreement or regulatory decree, then authoritarian countries would simply keep going. Given the incredible economic and military value of the technology, together with the lack of any meaningful enforcement mechanism, I don't see how we could possibly convince them to stop.

I do see a path to a *slight* moderation in AI development that is compatible with a realist view of geopolitics. That path involves slowing down the march of autocracies towards powerful AI for a few years by denying them the resources they need to build it, ⁴⁶ namely chips and semiconductor manufacturing equipment. This in turn gives democratic countries a buffer that they can "spend" on building powerful AI more carefully, with more attention to its risks, while still proceeding fast enough to comfortably beat the autocracies. The race

between AI companies within democracies can then be handled under the umbrella of a common legal framework, via a mixture of industry standards and regulation.

Contents

1. I'm sorry, Dave
2. A surprising and terrible empowerment
3. The odious apparatus
4. Player piano
5. Black seas of infinity
- Humanity's test

Anthropic has advocated very hard for this path, by pushing for chip export controls and judicious regulation of AI, but even these seemingly common-sense proposals have largely been rejected by policymakers in the United States (which is the country where it's most important to have them). There is so much money to be made with AI—literally trillions of dollars per year—that even the simplest measures are finding it difficult to overcome the political economy inherent in AI. This is the trap: AI is so powerful, such a glittering prize, that it is very difficult for human civilization to impose any restraints on it at all.

I can imagine, as Sagan did in *Contact*, that this same story plays out on thousands of worlds. A species gains sentience, learns to use tools, begins the exponential ascent of technology, faces the crises of industrialization and nuclear weapons, and if it survives those, confronts the hardest and final challenge when it learns how to shape sand into machines that think. Whether we survive that test and go on to build the beautiful society described in *Machines of Loving Grace*, or succumb to slavery and destruction, will depend on our character and our determination as a species, our spirit and our soul.

Despite the many obstacles, I believe humanity has the strength inside itself to pass this test. I am encouraged and inspired by the thousands of researchers who have devoted their careers to helping us understand and steer AI models, and to shaping the character and constitution of these models. I think there is now a good chance that those efforts bear fruit in time to matter. I am encouraged that at least some companies have stated they'll pay meaningful commercial costs to block their models from contributing to the threat of bioterrorism. I am encouraged that a few brave people have resisted the prevailing political winds and passed legislation that puts the first early seeds of sensible guardrails on AI systems. I am encouraged that the public understands that AI carries risks and wants those risks addressed. I am encouraged by the indomitable spirit of freedom around the world and the determination to resist tyranny wherever it occurs.

But we will need to step up our efforts if we want to succeed. The first step is for those closest to the technology to simply tell the truth about the situation humanity is in, which I have always tried to do; I'm doing so more explicitly and with greater urgency with this essay. The next step will be convincing the world's thinkers, policymakers, companies, and citizens of the imminence and overriding importance of this issue

Contents

1. I'm sorry, Dave
2. A surprising and terrible empowerment
3. The odious apparatus
4. Player piano
5. Black seas of infinity
- Humanity's test

—that it is worth expending thought and political capital on this in comparison to the thousands of other issues that dominate the news every day. Then there will be a time for courage, for enough people to buck the prevailing trends and stand on principle, even in the face of threats to their economic interests and personal safety.

The years in front of us will be impossibly hard, asking more of us than we think we can give. But in my time as a researcher, leader, and citizen, I have seen enough courage and nobility to believe that we can win—that when put in the darkest circumstances, humanity has a way of gathering, seemingly at the last minute, the strength and wisdom needed to prevail. We have no time to lose.



I would like to thank Erik Brynjolfsson, Ben Buchanan, Mariano-Florentino Cuéllar, Allan Dafoe, Kevin Esvelt, Nick Beckstead, Richard Fontaine, Jim McClave, and very many of the staff at Anthropic for their helpful comments on drafts of this essay.

Footnotes

¹ This is symmetric to a point I made in *Machines of Loving Grace*, where I started by saying that AI's upsides shouldn't be thought of in terms of a prophecy of salvation, and that it's important to be concrete and grounded and to avoid grandiosity. Ultimately, prophecies of salvation and prophecies of doom are unhelpful for confronting the real world, for basically the same reasons. ↩

² Anthropic's goal is to remain consistent through such changes. When talking about AI risks was politically popular, Anthropic cautiously advocated for a judicious and evidence-based approach to these risks. Now that talking about AI risks is politically unpopular, Anthropic continues to cautiously advocate for a judicious and evidence-based approach to these risks. ↩

³ Over time, I have gained increasing confidence in the trajectory of AI and the likelihood that it will surpass human ability across the board, but some uncertainty still remains. ↩

⁴ Export controls for chips are a great example of this. They are simple and appear to mostly just work. ↩

Contents

1. I'm sorry, Dave
2. A surprising and terrible empowerment
3. The odious apparatus
4. Player piano
5. Black seas of infinity
- Humanity's test

⁵ And of course, the hunt for such evidence must be intellectually honest, such that it could also turn up evidence of a lack of danger. Transparency through model cards and other disclosures is an attempt at such an intellectually honest endeavor. ↵

⁶ Indeed, since writing *Machines of Loving Grace* in 2024, AI systems have become capable of doing tasks that take humans several hours, with METR recently assessing that Opus 4.5 can do about four human hours of work with 50% reliability. ↵

⁷ And to be clear, even if powerful AI is only 1–2 years away in a technical sense, many of its societal consequences, both positive and negative, may take a few years longer to occur. This is why I can simultaneously think that AI will disrupt 50% of *entry-level* white-collar jobs over 1–5 years, while also thinking we may have AI that is more capable than *everyone* in only 1–2 years. ↵

⁸ It is worth adding that the *public* (as compared to policymakers) does seem to be very concerned with AI risks. I think some of their focus is correct (i.e. AI job displacement), and some is misguided (such as concerns about water use of AI, which is not significant). This backlash gives me hope that a consensus around addressing risks is possible, but so far it has not yet been translated into policy changes, let alone effective or well-targeted policy changes. ↵

⁹ They can also, of course, manipulate (or simply pay) large numbers of humans into doing what they want in the physical world. ↵

¹⁰ I don't think this is a straw man: it's my understanding, for example, that Yann LeCun holds this position. ↵

¹¹ For example, see Section 5.5.2 (p. 63–66) of the Claude 4 system card. ↵

¹² There are also a number of other assumptions inherent in the simple model, which I won't discuss here. Broadly, they should make us less worried about the specific simple story of misaligned power-seeking, but also more worried about possible unpredictable behavior we haven't anticipated. ↵

¹³ *Ender's Game* describes a version of this involving humans rather than AI. ↵

¹⁴ For example, models may be told not to do various bad things, and also to obey humans, but may then observe that many humans do exactly those bad things! It's not clear how this contradiction would resolve (and a well-designed constitution should encourage the model

to handle these contradictions gracefully), but this type of dilemma is not so different from the supposedly “artificial” situations that we put AI models in during testing. ↵

Contents

1. I’m sorry, Dave
2. A surprising and terrible empowerment
3. The odious apparatus
4. Player piano
5. Black seas of infinity
- Humanity’s test

¹⁵ Incidentally, one consequence of the constitution being a natural-language document is that it is legible to the world, and that means it can be critiqued by anyone and compared to similar documents by other companies. It would be valuable to create a race to the top that not only encourages companies to release these documents, but encourages them to be good. ↵

¹⁶ There’s even a hypothesis about a deep unifying principle connecting the character-based approach from Constitutional AI to results from interpretability and alignment science. According to the hypothesis, the fundamental mechanisms driving Claude originally arose as ways for it to simulate characters in pretraining, such as predicting what the characters in a novel would say. This would suggest that a useful way to think about the constitution is more like a character description that the model uses to instantiate a consistent persona. It would also help us explain the “I must be a bad person” results I mentioned above (because the model is trying to *act as if* it’s a coherent character—in this case a bad one), and would suggest that interpretability methods should be able to discover “psychological traits” within models. Our researchers are working on ways to test this hypothesis. ↵

¹⁷ To be clear, monitoring is done in a privacy-preserving way. ↵

¹⁸ Even in our own experiments with what are essentially voluntarily imposed rules with our Responsible Scaling Policy, we have found over and over again that it’s very easy to end up being too rigid, by drawing lines that seem important *ex ante* but turn out to be silly in retrospect. It is just very easy to set rules about the wrong things when a technology is advancing rapidly. ↵

¹⁹ SB 53 and RAISE do not apply at all to companies with under \$500M in annual revenue. They only apply to larger, more established companies like Anthropic. ↵

²⁰ I originally read Joy’s essay 25 years ago, when it was written, and it had a profound impact on me. Then and now, I do see it as too pessimistic—I don’t think broad “relinquishment” of whole areas of technology, which Joy suggests, is the answer—but the issues it raises were surprisingly prescient, and Joy also writes with a deep sense of compassion and humanity that I admire. ↵

Contents

1. I'm sorry, Dave
2. A surprising and terrible empowerment
3. The odious apparatus
4. Player piano
5. Black seas of infinity
- Humanity's test

²¹ We do have to worry about state actors, now and in the future, and I discuss that in the next section. ↩

²² There is evidence that many terrorists are at least relatively well-educated, which might seem to contradict what I'm arguing here about a negative correlation between ability and motivation. But I think in actual fact they are compatible observations: if the ability threshold for a successful attack is high, then almost by definition those who *currently* succeed must have high ability, even if ability and motivation are negatively correlated. But in a world where the limitations on ability were removed (e.g., with future LLMs), I'd predict that a substantial population of people with the motivation to kill but lower ability would start to do so—just as we see for crimes that don't require much ability (like school shootings). ↩

²³ Aum Shinrikyo did try, however. The leader of Aum Shinrikyo, Seiichi Endo, had training in virology from Kyoto University, and attempted to produce both anthrax and ebola. However, as of 1995, even he lacked enough expertise and resources to succeed at this. The bar is now substantially lower, and LLMs could reduce it even further. ↩

²⁴ A bizarre phenomenon relating to mass murderers is that the style of murder they choose operates almost as a grotesque sort of fad. In the 1970s and 1980s, serial killers were very common, and new serial killers often copied the behavior of more established or famous serial killers. In the 1990s and 2000s, mass shootings became more common, while serial killers became less common. There is no technological change that triggered these patterns of behavior, it just appears that violent murderers were copying each others' behavior and the "popular" thing to copy changed. ↩

²⁵ Casual jailbreakers sometimes believe that they've compromised these classifiers when they get the model to output one specific piece of information, such as the genome sequence of a virus. But as I explained before, the threat model we are worried about involves step-by-step, interactive advice that extends over weeks or months about specific obscure steps in the bioweapons production process, and this is what our classifiers aim to defend against. (We often describe our research as looking for "universal" jailbreaks—ones that don't just work in one specific or narrow context, but broadly open up the model's behavior.) ↩

²⁶ Though we will continue to invest in work to make our classifiers more efficient, and it may make sense for companies to share advances

Contents

1. I'm sorry, Dave
2. A surprising and terrible empowerment
3. The odious apparatus
4. Player piano
5. Black seas of infinity
- Humanity's test

like these with one another. ↵

²⁷ Obviously, I do not think companies should have to disclose technical details about the specific steps in biological weapons production that they are blocking, and the transparency legislation that has been passed so far (SB 53 and RAISE) accounts for this issue. ↵

²⁸ Another related idea is “resilience markets” where the government encourages stockpiling of PPE, respirators, and other essential equipment needed to respond to a biological attack by promising ahead of time to pay a pre-agreed price for this equipment in an emergency. This incentivizes suppliers to stockpile such equipment without fear that the government will seize it without compensation. ↵

²⁹ Why am I more worried about large actors for seizing power, but small actors for causing destruction? Because the dynamics are different. Seizing power is about whether one actor can amass enough strength to overcome everyone else—thus we should worry about the most powerful actors and/or those closest to AI. Destruction, by contrast, can be wrought by those with little power if it is much harder to defend against than to cause. It is then a game of defending against the most *numerous* threats, which are likely to be smaller actors. ↵

³⁰ This might sound like it is in tension with my point that attack and defense may be more balanced with cyberattacks than with bioweapons, but my worry here is that if a country's AI is the most powerful in the world, then others will not be able to defend even if the technology itself has an intrinsic attack-defense balance. ↵

³¹ For example, in the United States this includes the fourth amendment and the Posse Comitatus Act. ↵

³² Also, to be clear, there are some arguments for building large datacenters in countries with varying governance structures, particularly if they are controlled by companies in democracies. Such buildouts could in principle help democracies compete better with the CCP, which is the greater threat. I also think such datacenters don't pose much risk unless they are very large. But on balance, I think caution is warranted when placing very large datacenters in countries where institutional safeguards and rule-of-law protections are less well-established. ↵

³³ This is, of course, also an argument for improving the security of the nuclear deterrent to make it more likely to be robust against powerful AI, and nuclear-armed democracies should do this. But we don't know what a powerful AI will be capable of or which defenses, if any, will

Contents

1. I'm sorry, Dave
2. A surprising and terrible empowerment
3. The odious apparatus
4. Player piano
5. Black seas of infinity
- Humanity's test

work against it, so we should not assume that these measures will necessarily solve the problem. ↵

³⁴ There is also the risk that even if the nuclear deterrent remains effective, an attacking country might decide to call our bluff—it's unclear whether we'd be willing to use nuclear weapons to defend against a drone swarm even if the drone swarm has a substantial risk of conquering us. Drone swarms might be a new thing that is less severe than nuclear attacks but more severe than conventional attacks. Alternatively, differing assessments of the effectiveness of the nuclear deterrent in the age of AI might alter the game theory of nuclear conflict in a destabilizing manner. ↵

³⁵ To be clear, I would believe it is the right strategy not to sell chips to China, even if the timeline to powerful AI were substantially longer. We cannot get the Chinese “addicted” to American chips—they are determined to develop their native chip industry one way or another. It will take them many years to do so, and all we are doing by selling them chips is giving them a big boost during that time. ↵

³⁶ To be clear, most of what is being used in Ukraine and Taiwan today are not *fully* autonomous weapons. These are coming, but not here today. ↵

³⁷ Our model card for Claude Opus 4.5, our most recent model, shows that Opus performs better on a performance engineering interview frequently given at Anthropic than any interviewee in the history of the company. ↵

³⁸ “Writing all of the code” and “doing the task of a software engineer end to end” are very different things, because software engineers do much more than just write code, including testing, dealing with environments, files, and installation, managing cloud compute deployments, iterating on products, and much more. ↵

³⁹ Computers are general in a sense, but are clearly incapable on their own of the vast majority of human cognitive abilities, even as they greatly exceed humans in a few areas (such as arithmetic). Of course, things built *on top of* computers, such as AI, are now capable of a wide range of cognitive abilities, which is what this essay is about. ↵

⁴⁰ To be clear, AI models do not have precisely the same profile of strengths and weaknesses as humans. But they are also advancing fairly uniformly along every dimension, such that having a spiky or uneven profile may not ultimately matter. ↵

Contents

1. I'm sorry, Dave
2. A surprising and terrible empowerment
3. The odious apparatus
4. Player piano
5. Black seas of infinity
- Humanity's test

⁴¹ Though there is debate among economists about this idea. ↩

⁴² Personal wealth is a “stock,” while GDP is a “flow,” so this isn’t a claim that Rockefeller owned 2% of the economic value in the United States. But it’s harder to measure the total wealth of a nation than the GDP, and people’s individual incomes vary a lot per year, so it’s hard to make a ratio in the same units. The ratio of the largest personal fortune to GDP, while not comparing apples to apples, is nevertheless a perfectly reasonable benchmark for extreme wealth concentration. ↩

⁴³ The total value of labor across the economy is \$60T/year, so \$3T/year would correspond to 5% of this. That amount could be earned by a company that supplied labor for 20% of the cost of humans and had 25% market share, even if the demand for labor did not expand (which it almost certainly would due to the lower cost). ↩

⁴⁴ To be clear, I do not think actual AI productivity is yet responsible for a substantial fraction of US economic growth. Rather, I think the datacenter spending represents growth caused by anticipatory investment that amounts to the market expecting *future* AI-driven economic growth and investing accordingly. ↩

⁴⁵ When we agree with the administration, we say so, and we look for points of agreement where mutually supported policies are genuinely good for the world. We are aiming to be honest brokers rather than backers or opponents of any given political party. ↩

⁴⁶ I don’t think anything more than a few years is possible: on longer timescales, they will build their own chips. ↩

[Back to top](#)

[Privacy policy](#)

EXHIBIT 5

Type text here

Stenographic Transcript
Before the

COMMITTEE ON
ARMED SERVICES

UNITED STATES SENATE

TO RECEIVE TESTIMONY ON ENTERPRISE SECURITY AND
INFORMATION TECHNOLOGY OPERATIONS OF DEPARTMENT
OF DEFENSE NETWORKS
AND SYSTEMS

Tuesday, March 24, 2026

Washington, D.C.

ALDERSON COURT REPORTING
1029 VERMONT AVE, NW
10TH FLOOR
WASHINGTON, DC 20005
(202) 289-2260

TO RECEIVE TESTIMONY ON ENTERPRISE SECURITY AND INFORMATION
TECHNOLOGY OPERATIONS OF DEPARTMENT OF DEFENSE NETWORKS
AND SYSTEMS

Tuesday, March 24, 2026

U.S. Senate

Subcommittee on Cybersecurity

Committee on Armed Services

Washington, D.C.

The committee met, pursuant to notice, at 2:30 p.m. in
Room SR-232A, Russell Senate Office Building, Hon. Mike
Rounds, chairman of the subcommittee, presiding.

Committee Members Present: Senators Rounds
[presiding], Rosen, and Reed.

1 OPENING STATEMENT OF HON. MIKE ROUNDS, U.S. SENATOR
2 FROM SOUTH DAKOTA

3 Senator Rounds: Good afternoon, and welcome to this
4 afternoon's Cybersecurity Subcommittee hearing on enterprise
5 security and information technology operations of Department
6 of Defense Networks and Systems. I want to begin by
7 thanking our witnesses for appearing today before this
8 subcommittee. Your testimony arrives at an inflection point
9 for how the Department of Defense fights, decides, and wins.

10 This hearing is fundamentally about warfighting. The
11 digital backbone of the Department of Defense is no longer a
12 support function. It is a weapon system. Our tanks, ships,
13 aircraft, and ground forces depend on connected, resilient,
14 and secure networks to operate at the speed and precision
15 across vast distances. Sensors all over the world will
16 connect to shooters across the theater in a vast kill web
17 that must operate faster and more efficiently than our
18 enemies' systems. This is especially true for future
19 conflicts, as we continue fielding software intensive
20 systems designed for decision advantage and speed of action.

21 While quantity may have a quality all of its own, the
22 ability to optimize our targeting and orient, decide and act
23 more quickly than the enemy will likely decide the outcome
24 of the next major war. The Department's cultural shifts



1 towards recognizing networks and information technology
2 infrastructure as a warfighting platform is one I agree
3 with.

4 Yet, I realize there is still a long way to go.
5 Software can be developed and deployed in minutes, hours, or
6 days. The bureaucratic processes governing how we acquire,
7 certify and field it does not match that timeline, at least
8 not yet. This reality is no longer merely an inconvenience.
9 It is a strategic liability. We have talked long enough
10 about solving these problems. This subcommittee wants to
11 see them resolved.

12 Compounding the urgency is the condition of the
13 infrastructure itself. Defense information technology
14 budgets have long served as the bill payer absorbing cuts to
15 fund near-term priorities. The result is a technical debt
16 problem of historic proportions in both hardware and
17 software. Our adversaries are not blind to this. Not
18 surprisingly, they are taking advantage of this blunder.
19 Every day we delay modernizing is a day we strengthen their
20 hand in the cyber domain.

21 As we modernize, speed without security is not an
22 improvement. We must be disciplined about not punting
23 today's problems to the future. That means patching what we
24 have while building with foundational security baked in.



1 This cannot be a government-only endeavor. The private
2 sector can move with speed and provide capabilities the
3 Department cannot match. But partnership only works when
4 our processes and requirements are transparent and
5 consistent. We owe our industry partners clarity and we owe
6 our warfighters results. We must ask what kind of compute
7 capacity, network resiliency, and data infrastructure our
8 warfighters require, not just today, but as artificial
9 intelligence is deployed at a scale across the force. If we
10 cannot answer that question with specificity, we cannot
11 build toward it.

12 Today, this subcommittee looks forward to hearing from
13 both of you on where the Department stands in addressing
14 these efforts. We want to understand how bureaucratic
15 barriers are being dismantled, and how our processes are
16 being made more welcoming to industry partners. It is
17 especially important for us to hear whether our networks
18 have the capacity, resilience, and performance our military
19 requires for modern warfare against a capable adversary.

20 Thank you again, and now I would like to recognize the
21 ranking member for her remarks. Senator Rosen.

22

23

24



1 STATEMENT OF HON. JACKY ROSEN, U.S. SENATOR FROM
2 NEVADA

3 Senator Rosen: Well, thank you, Chairman Rounds. Ms.
4 Davies, General Stanton, appreciate you being here. I want
5 to welcome you, and Ms. Davies, congratulations on your
6 recent confirmation and assumption of duties. And so, Ms.
7 Davies, as you get settled into this position, we'd like to
8 hear more about your priorities, get some of the ideas of
9 where you're going to be placing emphasis on funding from
10 fiscal year 2026 appropriations and any reconciliation
11 funding you may have received as well.

12 And so, without discussing specifics, it would also be
13 helpful to know where we might see significant budget swings
14 for programs so we're not surprised when the fiscal year
15 2027 President's budget request is released. That request
16 is already nearly 2 months late. It's not expected for a
17 few more weeks, and that leaves little time in our NDAA
18 process to delve into the details and it's critical for us
19 to be able to do that.

20 General Stanton, I want to welcome you back. Thank you
21 for your years of service. I hope you will also share your
22 thoughts on all of these topics as they relate to your role
23 as director for the Defense Information Systems Agency and
24 commander for the Department of Defense Cyber Defense



1 Command. As the organization charged with running DOD
2 networks, you have a unique operational perspective on how
3 policy decisions from the Department CIO are translated,
4 well, into action, but that also means that you have a
5 slightly different view as to the resourcing and workforce
6 needs. The implementation of technology for the warfighters
7 different maybe than the development, right? The impact of
8 technical debt on the ability of our department to
9 modernize. You can speak to that in practical sense.

10 So, we also have a long list of long-term issues we
11 want to make sure the Department is addressing so they don't
12 get lost in the immediate priorities of the day or the
13 specific priorities of this administration. For example,
14 implementation of the cybersecurity maturity model
15 certification process, improving the authority to operate,
16 process for software to be placed on the DOD networks, and
17 addressing the backlog of technical debt in our IT programs
18 to make sure our networks are more modern, more resilient,
19 and of course, more secure.

20 These are just a few of the areas where I think we have
21 common cause and desire to work together to improve the
22 Department for the long term. We want to work
23 collaboratively with you on these issues. However, the
24 Department has not been particularly timely or forthcoming



1 on information. So, for us to conduct our statutory
2 obligation to oversee the Department of Defense, requires us
3 to have open and honest conversations. There's definitely a
4 trust deficit right now, but I think it's fixable if the
5 Department can be more responsive to the requests of this
6 committee and to our members. So, I hope we can use this
7 hearing, both open and closed, to help us get started on the
8 right foot.

9 And with that, I'm going to turn it back over to
10 Chairman Rounds. Thank you.

11 Senator Rounds: Thank you, Senator Rosen. Today,
12 we're pleased to have our two panelists with us, Kirsten
13 Davies, who is the chief information officer at the
14 Department of Defense, and Lieutenant General Paul Stanton,
15 United States Army director, Defense Information Systems
16 Agency commander, the Department of Defense Cyber Defense
17 Command. We welcome both of you here, and we look forward
18 to your opening statements. Your written statements will be
19 a part of the record, but we would welcome your opening
20 statements.

21 And with this, Ms. Davies, would you like to begin?
22
23
24



1 STATEMENT OF HONORABLE KIRSTEN A. DAVIES, CHIEF
2 INFORMATION OFFICER, DEPARTMENT OF DEFENSE

3 Ms. Davies: Thank you. Good afternoon, Chairman,
4 Ranking Member, Senator Reed, thank you for the opportunity
5 to speak with you today about the Department's strategy to
6 transform technology and cybersecurity into a decisive
7 warfighting advantage.

8 Our focus is to enable data supremacy and decision
9 dominance on the contested battlefields of today and
10 tomorrow at the speed and scale our warfighters deserve. In
11 the latest initiative in Secretary Hegseth's drive for
12 efficiency and effectiveness, we are undertaking a bold
13 transformation of enterprise IT and cybersecurity program,
14 unifying these capabilities under the Department's Chief
15 Information Officer. Through this effort, we will eliminate
16 inefficient spending, reduce technical debt, accelerate
17 modernization, drive consistent and up leveled
18 cybersecurity, and unleash data and innovation from the core
19 to the edge across our joint forces.

20 Leveraging my oversight of DISA, the NSA's
21 Cybersecurity Directorate, and the Department's Cyber Crime
22 Center, we are working with the military services, joint
23 staff, combatant commands, and defense agencies across four
24 transformation pillars. I'll provide a few highlights here.



1 Under pillar 1, the Enduring Digital Foundation, we're
2 transforming our network infrastructure and communications
3 transport, which extend from undersea cables to terrestrial
4 fiber to advanced satellite capabilities, connecting
5 everything from the home front to the tactical edge. This
6 foundation supports every warfighting system and our global
7 installations. We're driving continual modernization,
8 expansion, and hardening, as well as broad 5G usage and data
9 center modernization. We're evolving our cloud strategy in
10 JWCC Next, and we're also leading a proactive approach to
11 spectrum management and advancing PNT efforts, ensuring
12 ready and resilient capabilities that enable American
13 warfighting dominance.

14 Under pillar 2, agile digital capabilities, reflecting
15 some of your comments, Senator Rounds, we are expanding our
16 and maturing digital offerings. We're accelerating delivery
17 of software and SaaS services, and standardizing data
18 architectures, streamlining our data flows. We're shifting
19 from slow legacy software development to modern agile
20 delivery, driving interoperability by design, and delivering
21 applications and analytics at the speed of relevance. We're
22 driving extensive defense business systems work, whether
23 modernizing or sunseting those systems to enable clean
24 audits and reduce wasteful spends. We're also deploying



1 mission partner environment as persistent, secure
2 environments where trusted partners can be rapidly
3 integrated.

4 Under pillar 3, cybersecurity for the warfighting
5 ecosystem, in alignment with President Trump's National
6 Security Strategy and the National Defense Strategy, we're
7 moving from checklist-driven compliance towards unified,
8 holistic, risk-based approach. We will emphasize automation
9 and dynamic and continuous monitoring.

10 We will drive risk reduction rather than burdensome
11 paperwork, focusing on anti-fragility and resilience through
12 a holistic blend of streamlined processes, advanced
13 technologies, and skilled people. We're refining the
14 authority to operate, process, and accelerating our
15 deployment of Zero Trust principles. We're also refining
16 our risk management process and reigniting the debate to
17 align with the Secretary's arsenal of freedom initiatives.

18 Finally, through pillar 4, skills and partnerships, we
19 recognize that people are our decisive edge. As part of our
20 transformation, we're reviewing IT and cybersecurity roles
21 to ensure clear responsibilities, accountability for
22 outcomes, and a bias for action. We're leveraging your
23 provisions in the fiscal year 2026 NDAA to enhance
24 recruitment and retention of cyber professionals and expand



1 competitive compensation.

2 We will be launching an expanded top tier certification
3 program in partnership with industry and academia, which
4 will offer upskilling and reskilling to our warfighters,
5 from new recruits to seasoned service members. And because
6 we do not fight alone, we're doubling down to influence the
7 digital transformation efforts of our allies and partners,
8 which will better enable all of coalition force readiness.

9 As we advance this bold strategy, thank you for your
10 continued interest in and support of IT and cyber security,
11 and for the resources you provide to protect national
12 security. The race for data superiority and decision
13 dominance is won or lost every day, and this strategy is a
14 key part of how we, together, ensure the technology race is
15 won by the United States. Together, we will ensure the
16 resilience, readiness, and lethality of America's
17 warfighters across every domain.

18 I look forward to your questions.

19 [The prepared statement of Ms. Davies follows:]
20
21
22
23
24



1 Senator Rounds: Thank you, Ms. Davies. Lieutenant
2 General Stanton.

3
4
5
6
7
8
9
10
11
12
13
14
15
16
17
18
19
20
21
22
23
24



1 STATEMENT OF LIEUTENANT GENERAL PAUL T. STANTON, USA,
2 DIRECTOR, DEFENSE INFORMATION SYSTEMS AGENCY/ COMMANDER,
3 DEPARTMENT OF DEFENSE CYBER DEFENSE COMMAND

4 General Stanton: Chairman, Ranking Member Rosen,
5 Senator Reed, thank you for the privilege of appearing
6 before you today to explain the fundamental shift we are
7 making to ensure that our weapons system, the Department of
8 War Information Network, provides our warfighters with
9 decision advantage. I'm honored to represent the highly
10 skilled and dedicated professionals of the Department of
11 War, Cyber Defense Command, and the Defense Information
12 Systems Agency that design, build, secure, operate, and
13 defend our environment.

14 We must deliver a secure, standardized, resilient, and
15 efficient architecture that supports combatant commands.
16 Combatant commands execute warfighting. Nested within the
17 Department of War CIO's vision, the Defense Information
18 Systems Agency has a responsibility to provide functionally
19 relevant capability that aligns with the time and tempo of
20 the warfighter's mission. As the Department of War Cyber
21 Defense Command, we have an added responsibility to ensure
22 that our systems and data are properly defended against
23 continuous and sophisticated attacks.

24 Combining these two responsibilities, and said simply,

1 we must get the right data to the right place at the right
2 time, such that our commanders make better and faster
3 decisions than our enemies. We are a sub unified command
4 and a Department of War Combat Support Agency. Our mission,
5 and therefore our culture is to support warfighting. We are
6 fully engaged to move and maneuver the network and our data
7 according to the changing conditions of the operating
8 environment. We design, build, secure, operate, and defend
9 in lockstep with commanders at echelon. We campaign to
10 execute our missions and defeat our adversaries. We present
11 and defend the architecture so that commanders can fight.
12 We are doing so right now in Operation Epic Fury.

13 But we cannot rest. We must transform ourselves and
14 the means by which we support to leverage the most modern
15 and effective technology. We are in a perpetual state of
16 continuous modernization; new solutions, artificial
17 intelligence, commercial SATCOM, mobile data centers. They
18 emerge at industry's pace, and we must integrate them into
19 our architecture and missions at speed.

20 Further, we must be prepared to fight alongside our
21 partners, sharing data across warfighting functions within
22 decision cycles. The coalition information environment, as
23 recently prototyped during an exercise in the Indo-Pacific
24 theater, is a cornerstone of our approach. We know that our



1 adversaries, our enemies, are watching us and will certainly
2 attempt to delay or degrade our decisions, and we will
3 defeat them.

4 We are developing solutions that are secure by design,
5 incorporating perimeter defenses that defeat known attack
6 vectors and employing Zero Trust to detect, bound, and
7 defeat new and novel tradecraft. Our internet access point
8 and our cloud-based internet isolation warfighting systems
9 continuously adapt to new threats at our boundary. Our
10 Thunderdome implementation of Zero Trust is proven,
11 expanding rapidly as we transition defense agencies and
12 field activities into DODNet. We will use these tools and
13 the associated data in an informed, productive, efficient,
14 and speedy manner, automating our defenses and employing
15 human tradecraft for advanced analysis.

16 We prioritize our defenses on what matters. In order
17 to preserve decision space for commanders, we must defend
18 the critical systems upon which they are dependent. Our
19 approach employs a mission thread defense that we nest and
20 plan amongst commanders and their staffs for
21 synchronization. We align our cyber defenses to how their
22 systems employ and move data, ensuring that they have
23 confidence in the data upon which they make their decisions.
24 Our defenses and our strategic goals require optimization.



1 We have to see ourselves holistically. We have made
2 and continue to make significant progress in sensing,
3 logging, aggregating, and analyzing our data accordingly.
4 Our data analytics support cell is employing our common data
5 analytics platform to continuously run queries and analytics
6 that optimize performance and support defensive operations.
7 CDAO feeds into USCYBERCOM's joint cyber warfighting
8 architecture for enrichment with classified intelligence,
9 and coordination with offensive cyber forces for speed and
10 lethality.

11 Our success is the innovation, talent, and motivation
12 of our workforce. Readiness is a requirement. A trained
13 and ready force has confidence to act with disciplined
14 initiative. Demonstrated confidence leads to trust to make
15 decisions at speed. When we combine our transformational
16 architecture with a talented and trained workforce, we are
17 postured to meet our requirements. This is an imperative.
18 The effectiveness of the DoWIN is inextricably linked to our
19 missions and our Nation's defense. With the continued
20 support of the committee, we will preserve the decisive
21 advantage.

22 We look forward to your questions. Thank you.

23 [The prepared statement of General Stanton follows:]

24



1 Senator Rounds: Thank you, General Stanton.

2 Normally, we would begin with 5-minute rounds, and I
3 would start, Senator Rosen, the ranking member, would be
4 second, and then we'd move back and forth. But we also have
5 the ranking member here. And if you would like to --

6 Senator Reed: No sir. Regular order --

7 Senator Rounds: Okay. Regular order it is.

8 Senator Rosen: You heard it from the Ranking Member.

9 Senator Rounds: Very good. Well, then I will begin.
10 Ms. Davies and General Stanton, the ability of our
11 warfighters to operate in a contested or a degraded
12 environment is directly tied to how resilient and modern
13 those networks are. Where does the Department stand on
14 network modernization, and are our warfighters confident
15 they can operate if those networks are attacked or denied?
16 Ms. Davies.

17 Ms. Davies: Thank you, Senator Brown, for that great
18 question. I'll allow General Stanton to get into a few of
19 the details. But as I reflected on pillar 1 of our
20 transformation strategy, this is active work that we are
21 doing right now that DISA has been conducting for quite some
22 time under General Stanton's leadership and previous
23 leadership as well. This is a key area of focus for us.

24 Our forces are globally located. Our joint forces are



1 globally located. We are currently, obviously, in a mission
2 right now with partner forces as well. And the resiliency
3 and the efficacy of the traffic of our network is quite
4 critical to that. It's a key focus area for us, Senator.

5 Senator Rounds: General Stanton.

6 General Stanton: Senator, thank you for the question.
7 And with great support from Congress, we have an initiative
8 that we reference as design, security, and resiliency. So,
9 the Defense Information System Network, where we focus on
10 undersea cables, increased bandwidth for terrestrial fiber,
11 multimodal satellite communications capabilities we refer to
12 as an agnostic peering gateway that allows us to communicate
13 over military SATCOM waveforms, but also via commercial
14 SATCOM capabilities.

15 As our forces move into theater, as they currently
16 reside in theater, we have a primary alternate contingency
17 and emergency plans put into place. We're never single
18 threaded on any capability as we enter into the fight such
19 that if we suffer degradation, we have fallback
20 capabilities. We're seeing that in spades currently,
21 operating across terrestrial space based and undersea
22 capabilities.

23 Senator Rounds: So, recognizing that we're in an
24 unclassified environment, we'll go into a classified



1 environment when this when this meeting is done,
2 specifically, I think what you're indicating is there's a
3 couple of different areas where we may have, some
4 challenging communications problems. You mentioned,
5 undersea cables, you mentioned space-based assets and so
6 forth. Are those perhaps the most challenging that we're
7 going to face that we can talk about in this environment
8 today?

9 General Stanton: Well, Senator, I think it's the
10 combination and the fact that as our warfighting formations
11 are outfitted with capability. We never isolate down to a
12 single mode of transport. We ensure that we have the
13 ability to route terrestrially. We hit two peering points
14 such that we can leverage undersea cables. We're never
15 bounded by a single undersea cable. We always have a plan
16 to route around or have an alternate path.

17 And then, the proliferation of space-based assets,
18 specifically in the commercial world, is really game
19 changing technology to give us leap over capability if and
20 when we do suffer a degradation.

21 Senator Rounds: It's an interesting lead in on it.
22 Then for us to talk a little bit about the reason why we no
23 longer talk about a kill chain. We talk about a kill web,
24 and that is because we have multiple avenues to move from a



1 spotting system back into where you actually have the
2 ability to trigger a weapon. So, multiple ways to get the
3 communications from point A to point B, not simply one line.
4 Fair way of looking at it?

5 General Stanton: Precisely. Yes, sir.

6 Senator Rounds: Ms. Davies, many smaller Defense
7 Industrial Base companies lack the internal security
8 capability to defend themselves against a sophisticated
9 state actor. What is the Department doing to reach that
10 tier of the industrial base, and is voluntary participation
11 in government support programs getting us there?

12 Ms. Davies: Senator Rounds, this is a key focus area
13 for me as well. I think we have focused largely on
14 confidentiality of data in the past. I know that there has
15 been some burdensome requirements that have been placed on
16 large and small businesses alike. In the new
17 transformation, one of our key pillars is going to be
18 working directly with the Defense Industrial Base. Their
19 resiliency is our resiliency. Their security is our
20 security.

21 But it needs to make sense. We've heard Secretary
22 Hegseth talk about reducing the burdens to entrance,
23 allowing entrance, new entrance for smaller companies as
24 well. This is a key focus area for us, whether it's



1 providing guidance, providing principles for them to follow.
2 It's coming alongside them and partnering them, but it's
3 also tailoring these requirements so that they are effective
4 for the arsenal of freedom that we are driving.

5 Senator Rounds: Thank you. My time has expired.
6 Ranking Member Rosen.

7 Senator Rosen: Thank you, Chairman Rounds. I'm going
8 to just say something and I'll ask for more details in the
9 classified briefing. Just building on what Senator Rounds
10 said, I want to hear a little bit about the lessons you
11 learned from the recent military operations in Venezuela and
12 Iran that relate to the DOD networks in terms of showing us
13 a little bit of stress testing in the real world, right?
14 We're in conflicts in both places, and what we've learned
15 about what we might be changing for future protracted
16 conflicts. So, we'll save that one. Just put that out
17 there.

18 General Stanton, I want to talk a little bit about IT
19 support during a war. So, the Defense Information Systems
20 Agency, you're a combat support agency for the Department of
21 Defense. So, I'm hoping that you can explain for all of us
22 what that means, practically. Given our current posture in
23 the Middle East, what does your agency do during war time
24 that could be different than what it does during peace time,



1 if you could elaborate on that?

2 General Stanton: Yes, ma'am. Absolutely. Thank you
3 for the question. So, we are at war, and we're executing
4 Operation Epic Fury currently. Which means that every day
5 inside of our operations center, the Defense Information
6 Systems Agency, and Cyber Defense Command, get together to
7 ascertain what has transpired in the context of the network,
8 what assets are still up and running, what assets need to be
9 resolved? How do we route around problems? How do we
10 dynamically solve problems with emergent technology and do
11 so rapidly.

12 From a DISA perspective, a lot has to do with network
13 transport. And so, it -- where are the terminals located?
14 How do we get them to the right spot? Do we need to lease a
15 new circuit on the fly in order to ensure that critical data
16 gets from point A to point B? These problems present
17 themselves in real time, and inside of our ops center we are
18 dynamically solving and developing resolution.

19 Senator Rosen: Thank you. I'm going to move on to
20 you, Ms. Davies, because we want to talk a little bit about
21 the Joint Warfighting Cloud Contract. And the Joint
22 Warfighting -- it's a mouthful. The Joint Warfighting Cloud
23 Contract -- do not try to say that quickly -- we'll just say
24 the JWCC, it's a little bit easier, is reaching a point soon



1 where it's going to be need to be recommitted, and there'll
2 be a need to be a replacement contract. And so, what
3 lessons has DOD learned from the JWCC that we might see next
4 in an updated JWCC, and how do you see AI impacting your
5 decisions?

6 Ms. Davies: Thank you for the question. It is a
7 mouthful, isn't it?

8 Senator Rosen: It is, it is, yes.

9 Ms. Davies: We are expanding JWCC into a unified cloud
10 marketplace, integrating additional providers, embedding
11 financial operations, automation, and multi-cloud management
12 to enable enterprise-wide cost control and interoperability.
13 I can speak from being new in the role and seeing that there
14 are contracts everywhere, and different points of
15 authorization with different cloud that's happening.

16 One of the key areas that we need to be looking at from
17 a multi-pronged approach is, is this the most efficient way
18 to be driving cloud compute? It's not. We need to be
19 continuing on in the JWCC Next. Is it the most -- is it the
20 best way to see the spend. It is not. So, JWCC Next is
21 going to provide us that financial transparency that's
22 there, and it's also going to provide General Stanton and
23 his team the ability to do better defense across all this,
24 because we're going to know where all of the cloud compute



1 is, and that's key for us in asset identification and asset
2 security.

3 Senator Rosen: Thank you.

4 And then I'm going to continue with you, Ms. Davies,
5 because I want to talk about the enterprise chief
6 information officer collaboration. Right? So, DOD, you're
7 like -- you're saying you're just a vast conglomeration of
8 networks, clouds, operating cultures, systems, you name it.
9 It's difficult to craft a one-size-fits-all policy, although
10 you can set standards, and you can at least lay out the
11 templates for it. You can map out where everything is,
12 essentially, but in my view, there are benefits to having
13 each of the military departments and defense agencies having
14 their own CIOs so that they can tailor technology and
15 policies to the needs of the various organizations.

16 So, could you describe for us your relationship with
17 your CIO counterparts in the military services and defense
18 agencies, and are there any lessons, helpful or otherwise,
19 you might have picked up in your time so far?

20 Ms. Davies: Some great lessons indeed. I'm holding
21 regular meetings with my military department counterparts as
22 well as the DAFA counterparts. I'm learning where the
23 operational efficiencies are, where the centers of
24 excellence and expertise are, also where the gaps are. So,



1 I think there's varying levels of competencies, varying
2 levels of operational cadence that are there. And one of
3 the things that we will be getting after with this new
4 strategy is to take a hold of those rising tides, raise all
5 ships to make sure that we're all pointing in the same
6 direction and focused on operational excellence in cyber
7 defense.

8 Senator Rosen: I yield.

9 Senator Rounds: Senator Reed.

10 Senator Reed: Well, thank you, Mr. Chairman, Madam
11 Ranking Member. I thank the witnesses not only for being
12 here today, but for your dedication to our warfighters.
13 Thank you.

14 Ms. Davis, I'm sure you're aware of the recent decision
15 by the secretary to designate Anthropic as a supply chain
16 risk. Indeed, you, yourself, signed out a memo on March 6
17 directing the removal of Anthropic from DOD systems within
18 180 days. However, the committee still has not heard the
19 rationale from the Department about why the designation was
20 made, nor received a full notification, which is required
21 under law by Section 3252 of Title 10.

22 And this notification requires a summary of the risk
23 assessment and a summary of the basis for the determination,
24 including what less intrusive measures were considered, and



1 why they were not reasonably available to reduce supply
2 chain risk. We have not received that information, yet, it
3 is required under the law. So, first, are you aware of the
4 actual reason in designating Anthropic a supply chain risk?

5 Ms. Davies: Senator, thank you for the question. I
6 was involved, as many of my counterparts were, in this
7 collaborative decision-making process that followed the
8 regulatory requirements.

9 Senator Reed: Well, why was it done?

10 Ms. Davies: Sir, we're in active litigation right now,
11 so I won't go into the details of it. We have reached out
12 and offered a briefing for your offices into the depths of
13 it. I will say that there are some -- the filing in the
14 California court is available. We have made that available,
15 but it was only available to us this morning. So, we did
16 provide that over with the risk analysis. That was a part
17 of it.

18 Senator Reed: It's just interesting that you would
19 file required documentation for the California court before
20 complying with the law, and filing it, and sending it to us.
21 I don't believe it's been sent to us officially or
22 unofficially. Why haven't you, the Department, complied
23 with the law.

24 Ms. Davies: Senator, I'm aware that the -- our



1 colleagues in Legislative Affairs have followed the
2 regulation of what they were in -- what they were supposed
3 to provide. I do know that we followed all of the steps of
4 the regulatory requirement of the Title 10, 3252.

5 Senator Reed: Well, I don't think we've received it.
6 We have not received it.

7 Senator Rounds: Just in checking with staff, I do not
8 believe that we have received it at this time.

9 Senator Reed: Thank you, Mr. Chairman.

10 But as you pointed out, a California court has received
11 it because of the litigation. Let me just -- one additional
12 question is, are you aware of any estimates that were made
13 as to the potential cost impact on DOD uses for removing and
14 replacing Anthropic from DOD systems, or the cost to replace
15 Anthropic with another large language model.

16 Ms. Davies: Senator, I'm aware of the risk analysis
17 that was conducted as a part of that, and I'm aware that we
18 have also constructed our data architectures to be able to
19 be interoperable with a variety of different AI
20 capabilities. And so, the deep assessment of replacement of
21 that, I'm unfamiliar with right here, I can take that away
22 as an action for you.

23 [The information referred to follows:]

24 [SUBCOMMITTEE INSERT]



1 Senator Reed: It would be appreciated because the idea
2 of the scale and the magnitude of the disruption would be
3 helpful. Further, it's my understanding that Anthropic's
4 Claude system is being used today in Iran in our military
5 operations. Is that true?

6 Ms. Davies: Without going into the details in this
7 forum, Senator, the use of the system is active right now.
8 This is also why we provided for a measure of time we felt
9 was reasonable, as well as an exception process for removal
10 of the Anthropic systems.

11 Senator Reed: It just seems odd that you would
12 continue to use a system which you determined to be a supply
13 chain risk. Does that strike you as odd?

14 Ms. Davies: Senator, at no time, in any way, will we
15 interfere with the success, the lethality, and the
16 resilience of our warfighters. And for that reason, we've
17 provided what we feel is a reasonable amount of time for
18 those systems to be replaced. We can -- I can also say that
19 according to -- you know, with President Trump's great
20 leadership, we have a number of technology companies that
21 have come to the table wanting to do business with us as the
22 Department of War and across the U.S Government. So, we
23 know that we've architected this appropriately in order to
24 use competitive advantage as well.



1 Senator Reed: Thank you very much, Ms. Davis.

2 General, thank you.

3 General Stanton: Sure.

4 Senator Rounds: Thank you, Senator Reed.

5 Let me just follow-up on that for just briefly here.

6 My understanding is that there has been a 180-day
7 notification with regard to Anthropic. I presume, and you
8 can correct me if I'm wrong, but I presume that there is
9 additional time frame here in which there is the possibility
10 of additional negotiations that can occur during that time
11 period, recognizing just what a significant change this
12 would be to the Department with the reliance right now on
13 the Anthropic product at this time. Fair enough to say?

14 Ms. Davies: I'm not sure what part of the question to
15 answer for, Senator Rounds. Let me let me try to unpack
16 that for you. We have architected our data insomuch as we
17 can deploy multiple types of AI across our data. That's
18 something that the -- General Stanton and the DISA
19 colleagues have been very careful about. Number one.

20 Number two, we have provided what we feel is an
21 appropriate amount of time to remove the Anthropic systems
22 in accordance with the designation by the secretary of the
23 supply chain risk designation in and of itself. Does that
24 answer your question?



1 Senator Rounds: Yeah. Except that I think the other
2 entities that we are looking at, many of them have also
3 indicated that they have Anthropic within their systems as
4 well. And what I'm looking for is, is the possibility that
5 as this discussion goes on in a business-like manner, I'm
6 assuming it will be done in a business-like manner, that
7 there are opportunities for additional negotiations to
8 continue to occur?

9 Ms. Davies: Senator, I will defer that to my
10 colleagues in the legal department who are undergoing that
11 active litigation right now, and we will certainly bring a
12 report back to you.

13 Senator Rounds: That's fair. And I do think it'd be
14 fair to say that I think the committee as a whole, and I
15 can't speak for the chairman, but at least with regard to
16 the subcommittee, this is something that we have a real
17 interest in, and we will want to get in a classified setting
18 probably deeper into the details at some point when you're
19 prepared to share that with us -- with the appropriate
20 personnel.

21 Ms. Davies: Senator, we'll take that for action,
22 absolutely. Thank you.

23 Senator Rounds: Thank you.

24 Let me go on a little bit here. I'm just curious,



1 General Stanton, one item that we've talked about is
2 deterrence. And as we will use our offensive capabilities,
3 one of the reasons why you use offensive capabilities is to
4 deter future attacks, and to let people know that you know
5 who they are, we know where they are, and we do have access
6 to some very exquisite capabilities to stop them from
7 actually using kinetic activities or kinetic systems.

8 Can you talk a little bit, in this open session, just
9 so that the American public will understand, just kind of
10 some of the things that we have the ability to do now that
11 we've actually utilized some of them, and are -- the bad
12 guys know that what we can do? Can you talk just briefly
13 about that, just to share with the American public what
14 their taxpayer dollars are buying?

15 General Stanton: Yes, Senator. And I look forward to
16 having a more robust conversation in a classified setting.

17 Senator Rounds: I understand, but the public can't see
18 that. And like I said, I don't want to do any damage to our
19 ability to do it in the future. But I think it's fair for
20 deterrence's sake, to maybe talk a little bit about what our
21 capabilities are, if that is acceptable.

22 General Stanton: Yes, Senator. So, I think I'll
23 address it in two principal ways. The first is there's a
24 deterrent effect associated with cost imposition. If you



1 make it really hard for the enemy to attempt to achieve the
2 effects that the enemy intends, then it is a cost
3 imposition. The enemy has to spend more money, more time,
4 develop more resources, and apply it in ways that that the
5 enemy may not have been prepared. That's a cost imposition.

6 In addition, we have offensive cyber capabilities. And
7 we have offensive cyber capabilities that can, respond at
8 speed to evidence of adversarial activity beyond the bounds
9 of our U.S. networks. So, when we see operations in foreign
10 space as actioned by our cyberspace foreign adversaries, we
11 have the capabilities to deny them those resources.

12 Senator Rounds: Fair to say we can deny them the
13 ability to communicate in some cases?

14 General Stanton: Yes, Senator.

15 Senator Rounds: Fair to say that we can sometimes make
16 it so they can't see what they want to see on their systems
17 --

18 General Stanton: Yes, Senator.

19 Senator Rounds: -- to know what's going on in their in
20 their part of the world? Fair to say that we can make them
21 see things that maybe aren't even there today.

22 General Stanton: So, the ability to deny access to
23 systems, the ability to manipulate data to get inside the
24 decision cycle of the adversary, are all the art of the



1 possible in techniques developed in support of offensive
2 cyber operations.

3 Senator Rounds: All of which means that our young men
4 and women are then safer when they go in harm's way, because
5 we limit the adversary's ability to respond to our young men
6 and women who are on the battle front?

7 General Stanton: Unequivocally --

8 Senator Rounds: Thank you.

9 General Stanton: Unequivocally -- yes, Senator.

10 Senator Rounds: Thank you. Ranking Member Rosen.

11 Senator Rosen: Thank you.

12 I want to in the classified, I'm going to ask a little
13 bit more about how you've architected -- it's a new verb,
14 architected -- your data to feed into many different, I
15 would assume, large language models or the like. I think
16 that it's very interesting to me.

17 But for this open session, I want to talk about the
18 authority to operate process, because I'm encouraged by the
19 Department's continued support for improving cybersecurity
20 and supply chain risk management, of course, and you've made
21 progress towards reforming and accelerating acquisition,
22 testing, authorization of commercial software. But more, of
23 course, can always be done to streamline the authority to
24 operate, AT0, our single process into a single department-



1 wide accreditation for secure software providers,
2 streamlining the process.

3 So, can you talk about the status of your efforts to
4 streamline that process, of what actions you're taking, and
5 what plans your office may have to establish and encourage
6 reciprocity for ATOs between individual service branches and
7 department components. So, pulling all that back up to the
8 center.

9 Ms. Davies: Ranking Member, a great question. The ATO
10 process is part of the broader risk management framework, as
11 you are very familiar with. Right now, we have a very
12 static snapshot in time with regards to risk management, and
13 that needs to be moved to a much more dynamic framework and
14 process, which includes inheritance of assessments that are
15 conducted somewhere else across the Department on a piece of
16 software. That inheritance can then travel to a new
17 department that wants to leverage that piece of software,
18 that inheritance of all the testing and the scalability and
19 all of those types of things.

20 The ATO process as a piece of that also needs a reform.
21 We're finding that it's very difficult for people to
22 actually grab that inheritance over. It's very difficult to
23 understand the work of that risk management framework
24 because it's broken, it's fragmented, and it's static. So,



1 what we're doing is we're looking to do a lot more
2 automation across this, having dynamic repositories of this
3 information, and the testing in and of itself.

4 This work of the RMF framework, as well as the ATO
5 processes, will be sitting underneath the Department's chief
6 information security officer, who was presidentially
7 appointed just a few weeks ago. But he's right on top of it
8 and going to be working very actively with me to make sure
9 that we're reforming that appropriate to the risk of the
10 actual work that needs to be done.

11 Senator Rosen: Well, I understand what you're saying;
12 how important it is to be more dynamic, and sometimes less
13 static, but I'm also well aware of the vulnerabilities at
14 places when you are quickly dynamic without the proper
15 audits and controls over that as well because you never want
16 to sacrifice speed. And I'm not saying static is always the
17 way to go, but you have to be very careful about that
18 dynamic architecture as well because it can create
19 vulnerabilities, because moving at the speed of light, or
20 sound, or nano second, whatever you want, is -- has risk in
21 there as well. So, I hope that you're building in a good
22 audit process review of how that's working, so in that speed
23 --

24 Ms. Davies: Yes --



1 Senator Rosen: -- we will get to that.

2 Ms. Davies: -- and we're -- it's a great point. We're
3 going to be bringing in a lot of industry-good practices
4 across this as well, where we've learned to do a lot of the
5 automation of processes, less paperwork, more dynamic
6 checking across the software design lifecycle. We can do a
7 lot of code scanning, code reviews. Those are the types of
8 automation that will help us speed these things up while
9 we're compiling appropriate data repositories for that
10 constant risk checking.

11 Senator Rosen: Right. Because if you do it too fast,
12 once a piece of bad code gets in, it's already replicated
13 quickly across your system before you may have caught the
14 bug, so.

15 Ms. Davies: That's right. You understand the risk
16 process very much.

17 Senator Rosen: Yes, I think I do, but thank you.

18 I'm going to talk a little bit in my -- I just asked
19 you a question about artificial intelligence, the validity.
20 We talked about Anthropic. I know we're going to go talk
21 some more about this, the validity of their output critical
22 to the effectiveness of what we do. As we acquire and
23 deploy more systems, artificial intelligence systems,
24 warfighting -- you said we're in a war -- but it's going to



1 help our situational awareness, our decision-making, and we
2 must secure these systems and the data that powers them.
3 The data.

4 That's why I want to talk about how you architect your
5 data. Data is power if you're smart enough to analyze it
6 and use it. It's all about how you use it. The data is
7 key, and so it's important and critical to maintain the
8 trustworthiness, the integrity of all of that. Avoid any
9 corruption, malicious manipulation. And so, how are you
10 kind of -- as much as you can say here, what are you doing
11 to ensure that you're leveraging existing commercial
12 solutions, without making us vulnerable.

13 Ms. Davies: Yeah. Thank you, Ranking Member. In July
14 of 2025, my office published the DOW AI-Cybersecurity Risk
15 Management Tailoring Handbook. So, we provided some great
16 guidance across that. This is an ever-evolving framework or
17 competency, I would say. We've been tackling this in the
18 industry across the last, I'd say, 5 to 7 years, providing
19 appropriate guardrails, ethics, security across this,
20 looking at the weights of models, as well as hallucinations,
21 to try to reduce all of these factors that are there. We're
22 going to be continually evaluating this as we go through.

23 We have provided strong guardrails. We're doing risk
24 management assessing on a regular basis across this. And



1 so, this -- we will continue to have conversations around
2 this because it is an evolving category of software, if we
3 want to call it that.

4 Senator Rosen: Thank you.

5 Senator Rounds: I think at this time, we will conclude
6 the open portion of today's Cybersecurity Subcommittee
7 hearing, and as all of you know, we'll be reconvening here
8 in a few minutes. We've got a vote at 3:15 that's
9 scheduled. Matter of fact, three of them, but it'll give us
10 an opportunity to make our first vote.

11 We'd like to reconvene at 3:30 down in 217, in the
12 SCIF, for a classified portion of this. And then, for the
13 information of members who will not be joining us for the
14 closed briefing, questions for the record will be due to the
15 committee within 2 business days of the conclusion of the
16 hearing.

17 [The information referred to follows:]

18 [SUBCOMMITTEE INSERT]

19

20

21

22

23

24



1 Senator Rounds: And with that, I want to thank you for
2 this open session. We look forward to visiting with you
3 again very shortly, beginning at 3:30 in the closed session
4 in the SCIF.

5 And with that, the subcommittee meeting is adjourned.

6 [Whereupon, at 3:16 p.m., the hearing was adjourned.]
7
8
9
10
11
12
13
14
15
16
17
18
19
20
21
22
23
24



EXHIBIT 6

Anthropic

US National Security Agency using Anthropic's Mythos for cyber attacks

Arrangement comes as AI lab is locked in legal battle with Pentagon over Claude model



Anthropic has installed about half a dozen staff within the agency © Sebastien Bozon/AFP/Getty Images

Cristina Criddle in San Francisco and **Demetri Sevastopulo** in Washington

Published JUN 4 2026



Anthropic is helping the US National Security Agency deploy its powerful Mythos AI model for offensive cyber operations, embedding engineers inside the agency despite an ongoing legal battle with the Pentagon.

The San Francisco-based company had installed about half a dozen staff within the NSA as so-called forward-deployed engineers to guide the use of the technology and customise models for specific applications, two people familiar with the arrangement said.

It remains unclear whether Anthropic's engineers are assisting the NSA in active operations. However, one person close to the situation said Mythos would be useful for infiltrating the networks of nations such as China or Iran.

"The best way to build a good defence is to build a good attack," said a person close to Anthropic, who argued that adversaries were probably building their own [AI](#)-driven offensive technology. "If [Mythos] is not used to build attack agents, adversaries will find a way to do it."

The arrangement comes despite the Silicon Valley start-up's legal battle with the [defence department](#), which includes the NSA, over how its technology is used in warfighting.

[Anthropic](#) tried to restrict the US government's use of its Claude AI models for mass surveillance of American citizens and lethal autonomous drones. That led the Pentagon to label the AI lab a "supply-chain risk", a first for a US business. Anthropic has sued over the designation, which could force it to end contracts with organisations that work with the US military.

Since the dispute erupted, [Anthropic unveiled Claude Mythos](#), triggering concern among governments, financial institutions and IT companies over its ability to detect and exploit software vulnerabilities.

Earlier this week, Anthropic announced it would [distribute](#) Mythos to 150 organisations across 15 countries, vastly expanding access beyond the US and UK. It was initially made available only to a select number of US-based organisations when it launched in April.

The move comes after Anthropic [filed](#) for an initial public offering that could value the company at more than \$1tn, underlining the growing commercial and geopolitical significance of AI to national security.

Since Mythos's release, Anthropic rival OpenAI has released a model with similar capabilities. Experts have warned that advanced AI models enable hackers to exploit computer systems but can also serve as a powerful tool against adversaries.

Anthropic and the US defence department both declined to comment.

The company has worked closely with the US government on the wider rollout, while the Trump administration has embedded the model across agencies.

In February, [the FT reported](#) that the Pentagon was seeking to create AI-powered cyber tools to identify infrastructure targets in China as part of an effort to improve US capabilities in any future military conflict with Beijing.

On Tuesday, President Donald Trump signed an [executive order](#) outlining a voluntary framework in which AI companies can submit their new models for security reviews before they are publicly released.

The order directs federal agencies to develop methods to evaluate the cyber capabilities of AI systems and to establish an "AI cyber security clearinghouse" to share information and strengthen defences.

[Copyright](#) The Financial Times Limited 2026. All rights reserved.

EXHIBIT 7

White House and Anthropic Hold ‘Productive’ Meeting, Aiming for a Compromise

Friday’s meeting at the White House followed the introduction of Anthropic’s powerful new artificial intelligence model, Mythos, which U.S. officials believe could be critical for security.



By **Julian E. Barnes**, **Sheera Frenkel** and **Tyler Pager**

Julian E. Barnes and Tyler Pager reported from Washington, and Sheera Frenkel from San Francisco.

April 17, 2026

See more of our coverage in your search results.

Add The New York Times on Google ↗

Anthropic’s chief executive met with White House officials on Friday for discussions that both sides described as “productive,” as the Trump administration works to forge a compromise that would bring the artificial intelligence company’s technology back into the government, according to U.S. officials and others briefed on the matter.

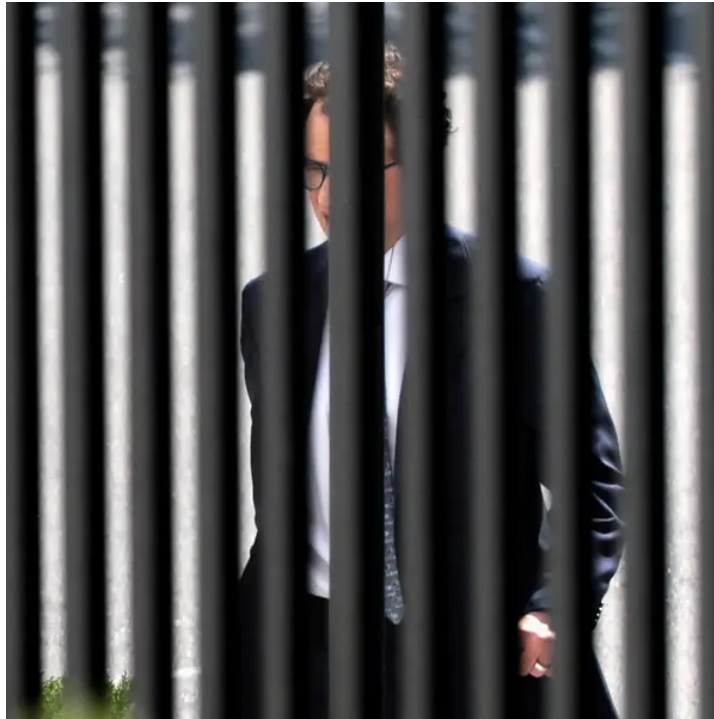
Dario Amodei, Anthropic’s chief executive, met with the White House chief of staff, Susie Wiles; Treasury Secretary Scott Bessent; and others, the people said.

The meeting was “both productive and constructive,” the White House said in a statement, with conversations about how to collaborate and address the challenges of A.I. while exploring “the balance between advancing innovation and ensuring safety.” Anthropic said in a statement that the discussions had touched on “key shared priorities” such as cybersecurity and A.I. safety.

Anthropic had been effectively cut off from working with the federal government after [battling with the Pentagon](#) this year. The two sides disagreed in negotiations over a \$200 million contract about the [use of A.I. in warfare](#), with the Pentagon later designating Anthropic a “[supply chain risk](#).”

Friday’s meeting was a potential first step to a deal. If officials reach a compromise with the company, it is likely to exclude the Pentagon, two officials said.

Some White House officials have argued that a fight with Anthropic is counterproductive and denies the United States some of the most powerful tech tools, officials said. This follows Anthropic’s unveiling on April 7 of a powerful new A.I. model, [Mythos](#), which is capable of identifying security vulnerabilities in software. Officials believe access to the model — which Anthropic has made available only on a limited basis — is critical to help protect government networks from cyberattacks.



Dario Amodei, Anthropic's chief executive, arriving for the meeting with officials at the White House on Friday. *Jessica Koscielnia/Reuters*

Some White House officials were also frustrated that other officials had failed to find a way to de-escalate the contract fight with Anthropic, especially given the potential for Mythos to wreak havoc on computer systems, officials said.

“We look forward to continuing this dialogue and will host similar discussions with other leading A.I. companies,” the White House statement said.

Anthropic added that the meeting had reflected its “ongoing commitment to engaging with the U.S. government on the development of responsible A.I.” and that it was “looking forward to continuing these discussions.”

The White House meeting [was reported earlier by Axios](#).

Neither the Pentagon nor Anthropic had shown much willingness to settle their dispute. During contract negotiations, Anthropic had sought assurances that its powerful A.I. models — which until recently were the only ones allowed on classified computers at the Defense Department — would not be used for commanding autonomous lethal weapons or surveilling Americans. In response, the Pentagon said no private contractor could tell it how to use the technology.

The fight came to a head on March 5, when Secretary of Defense Pete Hegseth labeled Anthropic a supply chain risk. The designation was previously used only for foreign companies that the government believed posed a risk to national security. Anthropic later [sued the U.S. government](#) over the designation in courts in California and Washington, D.C.



Secretary of Defense Pete Hegseth on Thursday. Anthropic and the Pentagon disagreed over the use of its technology in warfare earlier this year. **Pete Marovich for The New York Times**

Since then, Anthropic has told government officials that it is willing to provide access to Mythos to help agencies use the tool to find and fix potential vulnerabilities in software and computer networks.

At the Pentagon, engineers who work with Anthropic are petitioning the department to keep using the company's technology, according to two officials who have participated in meetings about A.I. technology this month. While the engineers were not averse to using new A.I. models from other companies, they did not want to be cut off from Anthropic's technology, which is used to analyze intelligence and handle sensitive data.

The engineers also urged the Pentagon to update to the newest Anthropic models available, the two officials said. Because the models are held on systems housed within the Pentagon, Anthropic cannot update them unless given access by defense officials, they said.

Senior Pentagon officials declined to talk about Anthropic, citing the ongoing lawsuits. The department is using an older version of Anthropic's Opus A.I. model than it would have had the dispute not occurred, current and former defense officials said.

Other officials said they were pressing ahead with bringing models from OpenAI and Google to Pentagon computers, including a more advanced version of Google's Gemini that will be online for military use in the coming days.

After Anthropic sued the U.S. government over the "supply chain risk" label, a judge in the U.S. District Court for the Northern District of California [temporarily stopped](#) the Pentagon from enforcing the designation in late March. In a separate ruling in the U.S. Court of Appeals for the District of Columbia Circuit on April 8, judges [denied Anthropic's request](#) to stop the Pentagon from labeling the company a supply chain risk.

Julian E. Barnes covers the U.S. intelligence agencies and international security matters for The Times. He has written about security issues for more than two decades.

Sheera Frenkel is a reporter based in the San Francisco Bay Area, covering the ways technology affects everyday lives with a focus on social media companies, including Facebook, Instagram, Twitter, TikTok, YouTube, Telegram and WhatsApp.

Tyler Pager is a White House correspondent for The Times, covering President Trump and his administration.

A version of this article appears in print on , Section B, Page 4 of the New York edition with the headline: Trump Team and Anthropic Try to Reach a Compromise