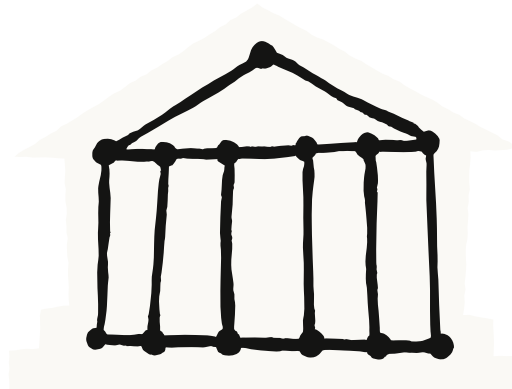


Announcements

A statement from Dario Amodei on Anthropic's commitment to American AI leadership

Oct 21, 2025 • 6 min read



A statement from Anthropic CEO Dario Amodei on Anthropic's commitment to advancing America's leadership in building powerful and beneficial AI.

Anthropic is built on a simple principle: AI should be a force for human progress, not peril. That means making products that are genuinely useful, speaking honestly about risks and benefits, and working with anyone serious about getting this right. I strongly agree with Vice President JD Vance's recent comments on AI

—particularly his point that we need to maximize applications that help people, like breakthroughs in medicine and disease prevention, while minimizing the harmful ones. This position is both wise and what the public overwhelmingly wants.

Anthropic is the fastest-growing software company in history, with revenue growing from a \$1B to \$7B run rate over the last nine months, and we've managed to do this while deploying AI thoughtfully and responsibly. There are products we will not build and risks we will not take, even if they would make money.

Our longstanding position is that managing the societal impacts of AI should be a matter of policy over politics. I fully believe that Anthropic, the administration, and leaders across the political spectrum want the same thing: to ensure that powerful AI technology benefits the American people and that America advances and secures its lead in AI development.

Despite our track record of communicating frequently and transparently about our positions, there has been a recent uptick in inaccurate claims about Anthropic's policy stances. Some are significant enough that they warrant setting the record straight.

Our alignment with the Trump administration on key areas of AI policy

- We work directly with the federal government in several ways. In July the Department of War awarded Anthropic a two-year, \$200 million agreement to prototype frontier AI capabilities that advance national security. We have partnered with the General Services Administration to offer Claude for Enterprise and Claude for Government for \$1 across the federal government. And Claude is deployed across classified networks through partners like Palantir and at Lawrence Livermore National Laboratory.
- Anthropic publicly praised President Trump's AI Action Plan. We have been supportive of the President's efforts to expand energy provision in the US in order to win the AI race, and I personally attended an AI and energy summit in Pennsylvania with President Trump, where he and I had a good conversation about US leadership in AI. Anthropic's Chief Product Officer attended a White House event where we joined a pledge to accelerate healthcare applications of AI, and our Head of External Affairs attended the

White House's AI Education Taskforce event to support their efforts to advance AI fluency for teachers.

- Every major AI company has hired policy experts from both parties and recent administrations—Anthropic is no different. We've hired Republicans and Democrats alike, and built an advisory council that includes senior former Trump administration officials. Anthropic makes hiring decisions based on candidates' expertise, integrity, and competence, not their political affiliations.
- We (and many other organizations) respectfully disagreed with a single proposed amendment in the One Big Beautiful Bill: the 10-year moratorium on state-level AI laws, which would have blocked any action without offering a federal alternative. That specific provision was voted down by Republicans and Democrats in a 99-1 vote in the Senate. Our longstanding position has been that a uniform federal approach is preferable to a patchwork of state laws. I proposed such a standard months ago and we're ready to work with both parties to make it happen.

Our preference for a national AI standard

- While we continue to advocate for a national standard, AI is moving so fast that we can't wait for Congress. **Skip to footer** We supported a carefully designed bill in California where most of America's leading AI labs are headquartered, including Anthropic. This bill, SB 53, requires the largest AI developers to make their frontier model safety protocols public and is written to exempt any company with an annual gross revenue below \$500M—therefore only applying to the very largest AI companies. Anthropic supported this exemption to protect startups and in fact proposed an early version of it.
- Some have suggested that we are somehow interested in harming the startup ecosystem. Startups are among our most important customers. We work with tens of thousands of startups and partner with hundreds of accelerators and VCs. Claude is powering an entirely new generation of AI-native companies. Damaging that ecosystem makes no sense for us.
- I've heard arguments that state AI regulation could slow down the US AI industry and hand China the lead. But the real risk to American AI leadership isn't a single state law that only applies to the largest companies—it's filling

the PRC's data centers with US chips they can't make themselves. We agree with leaders like Senators Tom Cotton and Josh Hawley that this would only help the Chinese Communist Party win the race to the AI frontier. We are the only frontier AI company to restrict the selling of AI services to PRC-controlled companies, forgoing significant short-term revenue to prevent fueling AI platforms and applications that would help the Chinese Communist Party's military and intelligence services.

Our progress on an AI industry-wide challenge: model bias

- Some have claimed that Anthropic's models are uniquely politically biased. This is not only unfounded but directly contradicted by the data. A January study from the Manhattan Institute, a conservative think tank, found Anthropic's main model (at the time, Claude Sonnet 3.5) to be less politically biased than models from most of the other major providers. Data from a Stanford study in May, on user perceptions of bias in AI models, shows no reason to single out Anthropic: many models from other providers were rated as more biased. The system cards for our latest models, Sonnet 4.5 and Haiku 4.5, show that we're making rapid progress towards our goal of political neutrality.
- As a broader point, no AI model from any provider is fully politically

AI

≡

outputs. Anyone can cherry-pick outputs from any model to make it appear slanted in a particular direction.

Anthropic is committed to constructive engagement on matters of public policy. When we agree, we say so. When we don't, we propose an alternative for consideration. We do this because we are a public benefit corporation with a mission to ensure that AI benefits everyone, and because we want to maintain America's lead in AI. Again, we believe we share those goals with the Trump administration, both sides of Congress, and the public. We are going to keep being honest and straightforward, and will stand up for the policies we believe are right. The stakes of this technology are too great for us to do otherwise.

In his recent remarks, the Vice President also said of AI, "Is it good or is it bad, or is it going to help us or going to hurt us? The answer is probably both, and we should be trying to maximize as much of the good and minimize as much of the