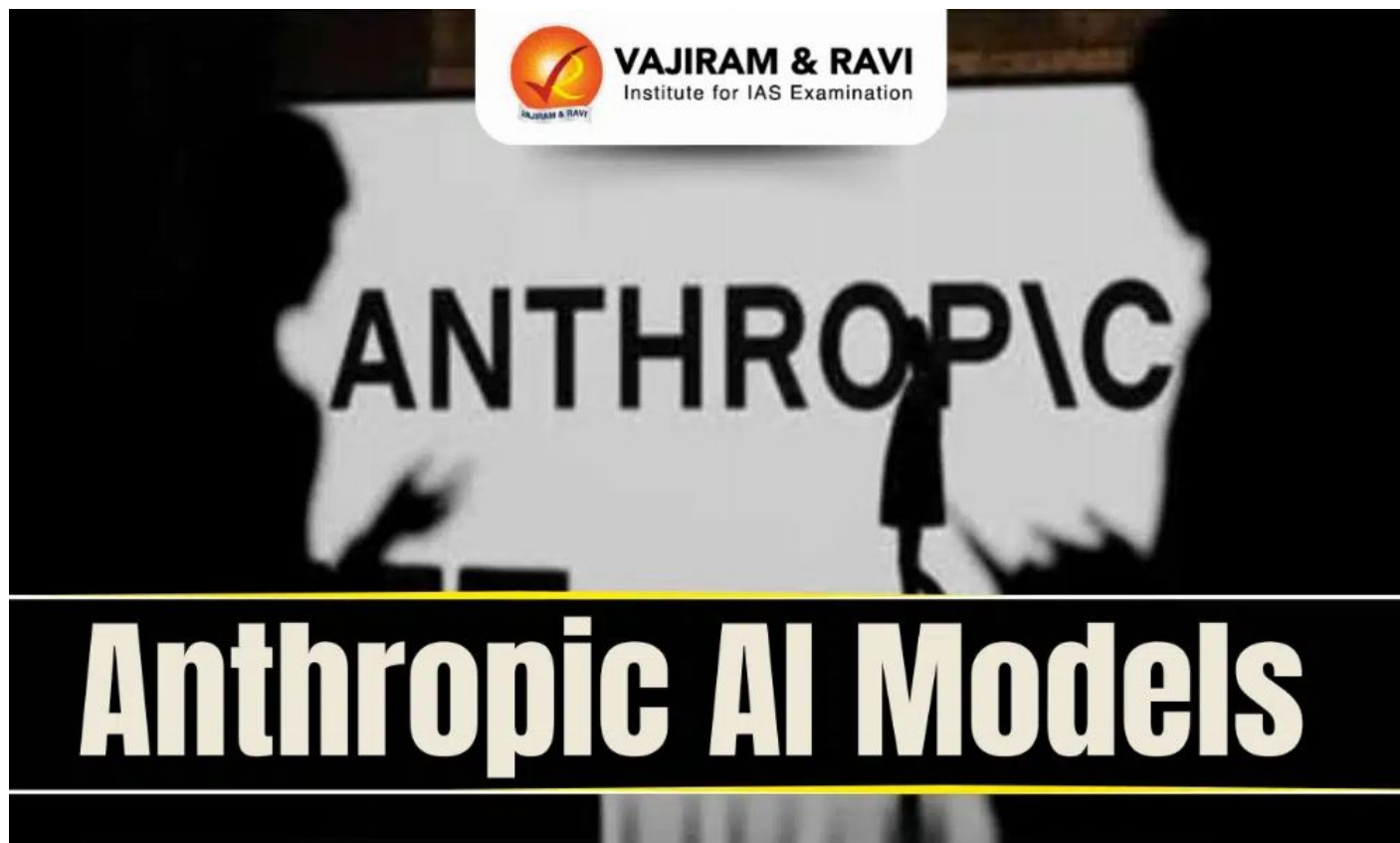


Anthropic AI Models: Why Foreign Access to Fable 5 and Mythos 5 Was Restricted

June 14, 2026 • vajiramandravi.com



Anthropic AI Models Latest News

- The US government has issued an **export control order** barring **all foreign nationals** — even those working inside Anthropic, a leading AI company — from accessing its two newest AI models, **Fable 5 and Mythos 5**.
- These were launched recently and are Anthropic's most advanced publicly available AI systems.
- The order cites **national security concerns**, but its broad scope has raised fresh questions about how governments can control access to AI technology.

What Are Export Controls, and Why Is This Different

- Export controls are rules that stop certain goods or technologies from being sent to other countries. Normally, they apply to **physical items** — for example, the US restricts the sale of advanced computer chips to China.

- This new order is different. For the first time, it controls how an AI software product can be used and by whom — not a physical product.
- This is being seen as a new and unusual use of export control law, applied to artificial intelligence for the first time in this manner.

Export Control Order Banning Access of AI Model

- The order blocks **all foreign nationals** from using Fable 5 and Mythos 5 — whether they are located inside the US or outside it. This includes Anthropic’s own foreign employees.
- In simple terms, only American citizens or entities can access these two models for now.
- As per the Anthropic, the government did not give a clear, written explanation. Based on informal communication, the company believes the **concern relates to a possible “jailbreak”** — a method to bypass an AI model’s safety restrictions.
 - This technique works by asking the AI model to read a piece of software code and identify and fix flaws in it.
- Anthropic has pushed back, stating that this capability is not unique to its models — other publicly available AI systems, including OpenAI’s GPT-5.5, can do the same thing.
- In fact, this kind of capability is used daily by cybersecurity professionals to find and fix software weaknesses.

Background: Anthropic’s Troubled Relationship with the US Government

- This is not the first friction between Anthropic and US authorities. In March 2026, the Pentagon labelled Anthropic a **“supply chain risk”** — a designation that could limit how much US government agencies can use the company’s AI products.
- This labelling followed a disagreement over military use of Anthropic’s AI (Claude).
 - Anthropic had insisted on certain safety conditions — specifically, that its AI should not be used for mass domestic surveillance or to help build fully autonomous weapons systems.
 - The company has reportedly indicated it may legally challenge this “supply chain risk” label.
- This export control order comes at a sensitive time — Anthropic is preparing for a public stock market listing in the US, and was recently valued at \$965 billion.

What Are Fable 5 and Mythos

- Anthropic has launched **Fable 5** and **Mythos 5**, its most advanced publicly available AI models.
- Both are built on the company’s new **“Mythos-class” architecture**, derived from the earlier **Mythos Preview** model, which was not publicly released due to concerns over potential misuse.

Why Was Mythos Preview Not Released

- Anthropic had previously claimed that Mythos Preview possessed the capability to identify severe vulnerabilities in major operating systems and web browsers, including some that reportedly remained undetected for years.

- The company feared that such powerful capabilities could be misused, particularly against critical infrastructure worldwide. As a result, it decided against a public release.
- Before the launch, Anthropic provided limited access to Mythos Preview through Project Glasswing.
- In India, the company reportedly discussed sharing the model with a small number of organizations.

Key Difference Between Mythos Preview and Fable 5/Mythos 5

- Although Fable 5 and Mythos 5 are based on the same underlying technology as Mythos Preview, they include **additional safety safeguards**.
 - Requests related to sensitive cybersecurity issues or biological threats are filtered.
 - When the system detects a potentially high-risk query, it automatically redirects the user to a less capable AI model (Claude Opus 4.8) instead of allowing the full Mythos-level system to respond.
- The company acknowledges that the safeguards may occasionally block harmless requests.

Conclusion

- This episode highlights a growing global concern.
- As AI models become powerful enough to find security flaws in critical software, governments are beginning to treat advanced AI itself as a controlled, sensitive technology — similar to how they treat nuclear technology, advanced weapons, or top-end semiconductor chips.
- For India, this is a reminder that building indigenous AI capability is no longer optional but a matter of strategic necessity.

Source: **IE** | **ToI**

Source: <https://vajiramandravi.com/current-affairs/anthropic-ai-models/>

© Vajiram & Ravi. For personal use only.