



Ceasefire, Not Peace Treaty: What the Lutnick Letter Means for AI Policy

On June 30, 2026, Commerce Secretary Howard Lutnick formally withdrew the export controls that had forced Anthropic's Claude Fable 5 and Mythos 5 models offline for 18 days. The immediate crisis has passed but the terms of the resolution, and what they signal about AI policy going forward, warrant careful scrutiny.

On June 12, the U.S. Bureau of Industry and Security (BIS) issued an unprecedented "is informed" letter invoking the Export Control Reform Act, requiring Anthropic to obtain a license before sharing Fable 5 or Mythos 5 with any foreign national – including its own employees. Anthropic was given roughly 90 minutes to comply with the direction, and – lacking any way to verify nationality in real time – suspended access to both models for all users.

The trigger was a jailbreak report. Amazon researchers had bypassed Fable 5's safeguards to identify software vulnerabilities and produce exploit demonstration code. Anthropic's testing confirmed that less capable models – including GPT-5.5 and Claude Opus 4.8 – could reproduce the same results, and more than 80 cybersecurity executives signed an open letter calling the threat overstated.

The Commitments Behind the Resolution

Lutnick's June 30 letter identifies three commitments Anthropic made to

- Proactively detect and address security risks;
- Work with the government on protocols and standards for current and future model releases; and
- Inform the government of malicious activity.

In its June 30 announcement, Anthropic detailed the operational specifics. The company trained an improved safety classifier that blocks the specific jailbreak technique in over 99% of cases. The Commerce's Center for AI Standards and Innovation independently tested the safeguards and agreed they are "extraordinarily strong."

Anthropic also committed to deeper government collaboration: pre-release access and evaluation for frontier models, rapid information sharing on jailbreaks and participation in the interagency vulnerability clearinghouse, dedicated compute for joint government research, and work toward a shared industry security standard. Additionally, Anthropic partnered with Amazon, Microsoft, Google, and other Glasswing partners to develop a consensus framework for scoring jailbreak severity – modeled on the Common Vulnerability Scoring System for software vulnerabilities.

A Conditional Reprieve: What the Lutnick Letter Says and What It Doesn't

The letter resolves the immediate standoff but explicitly preserves government discretion. Its final paragraph states that "Commerce reserves the right to reevaluate the decisions made in this letter and the necessity of reimposing a license requirement, should circumstances change or should Anthropic fail to adhere to its commitments." This is a conditional reprieve, not a permanent resolution. Anthropic operates under commitments whose specific terms remain undisclosed – and under the implicit threat that any perceived breach could trigger a repeat.

Notably, the letter was addressed to Anthropic's Chief Compute Officer Tom Brown – not CEO Dario Amodei, who had been a target of the administration – reinforcing the impression that the resolution was negotiated around the political friction point rather than through it.

Yet the letter does not provide defined criteria for reimposition. There is no process guarantees (notice, hearing, or appeal) and no industry-wide standard. The commitments are bilateral, leaving other frontier model providers in regulatory limbo. For downstream customers, the letter offers no direct protection – controls were imposed and lifted entirely through bilateral communications, with no notice to or input from commercial licensees whose businesses were disrupted.

The Emerging Regulatory Playbook

The Anthropic episode is a proof of concept for how the current administration may exercise AI regulatory power. The government has now demonstrated that it will use Export Control Reform Act of 2018 export control authority to pull a commercially deployed AI model offline over a jailbreak claim – even one that other models can replicate – setting a concrete and unprecedented precedent for software.

The tension between stated policy and actual practice is striking. The June 2 Executive Order (EO) explicitly disclaims any “mandatory governmental licensing, preclearance, or permitting requirement” for AI models. Yet the Anthropic episode shows that ad hoc directives backed by criminal penalties achieve functionally mandatory compliance – making the voluntary framework coercive in practice. Meanwhile, staged access is quietly becoming the norm: OpenAI’s GPT-5.6 was released only to approved customers after a government request, and the June 2 EO contemplates up to 30 days of pre-release government access for frontier models.

The next concrete governance milestone is the August 2026 deadline. The EO gives national security and cybersecurity officials 60 days to develop a classified benchmarking process designating “covered frontier models.” How those benchmarks are calibrated will determine how frequently export-control-style interventions recur. Yet even that deadline will not resolve the deeper problem: there is currently no clear regulatory home or transparent process for AI governance. Until a durable rulemaking process exists –grounded in statutory authority and subject to notice-and-comment – companies will continue to face the structural uncertainty that ad hoc directives create.

Practical Takeaways for In-House Counsel

The Anthropic episode exposes a category of regulatory risk that most AI-related contracts do not adequately address. Counsel should consider the following:

1. *Force majeure and regulatory-action clauses need updating.* Traditional force majeure language contemplates natural disasters, war, and sometimes government action. Contracts for AI-dependent services should specifically address sudden model unavailability due to government action, including

whether and how quickly the provider must migrate to an alternative model, and how performance obligations are tolled or excused during the disruption.

2. *Termination and suspension rights require precision.* The 18-day suspension fell in a gap between short-term service outages (typically covered by SLA credits) and long-term unavailability (which might trigger termination for cause). Counsel should draft suspension thresholds that distinguish between technical downtime and regulatory-forced unavailability, with escalating remedies – notice obligations, fee abatement, migration assistance, and ultimately termination rights – calibrated to the duration and cause of the disruption.
3. *Vendor diversification covenants may become standard.* A customer whose product depends entirely on a single frontier AI model has no fallback when that model is pulled. In-house counsel negotiating AI service agreements should consider requiring vendors to deliver a remediation plan and demonstrate substitution readiness.
4. *Regulatory risk allocation clauses deserve explicit treatment.* The Lutnick letter's reservation of rights means the regulatory risk has not been eliminated – it has been deferred. Contracts should allocate the risk of future export control actions, including: (a) which party bears the cost of compliance with new government-imposed safeguards; (b) whether government-mandated changes to model capabilities (e.g., expanded safety classifiers that increase false positives) constitute a material change in service; and (c) indemnification for losses flowing from a provider's failure to maintain its regulatory commitments.
5. *Representations regarding regulatory status are now material.* AI vendors should be asked to represent that their models are not subject to pending or threatened export control actions, "is informed" letters, or other BIS directives – and to covenant that they will promptly notify customers if such actions are initiated. The absence of any customer-notification obligation in the Lutnick framework makes contractual notice rights essential.

The Bottom Line: Stability Requires Rules, Not Directives

The Lutnick letter is a ceasefire, not a peace treaty. It resolves a bilateral dispute while leaving the underlying policy vacuum intact. Companies should anticipate that pre-release government evaluation will become standard, that export control authority will remain an emergency lever, and that the real regulatory framework will

be shaped by the August 2026 benchmarking process and whatever emerges from Congress.

The fundamental question remains: will AI policy be made through transparent, accountable rulemaking, or by ad hoc directives? The Anthropic resolution suggests the government is willing to negotiate outcomes – but only after acting first and asking questions later. For an industry that needs stability to invest and compete, that sequence matters.

This blog was drafted by [Diana Shutzer](#) and [Nicholas Solosky](#), leaders of the Spencer Fane Government Contracts team. For more information, visit spencerfane.com.

Click [here](#) to subscribe to Spencer Fane communications to ensure you receive timely updates like this directly in your inbox.