



By: Tomorrow's Affairs Staff

Who decides how America uses artificial intelligence in war?



On 27 August, **US federal judge** Rita F. Lin overturned the Donald Trump administration's decision to designate Anthropic as a supply-chain risk and permanently banned the enforcement of measures based on that decision.

In a 59-page **ruling**, she concluded that the Pentagon had used an instrument intended to protect national security without a sufficient legal basis, that Anthropic was exposed to retaliation for its publicly stated position, and that the company had not been afforded the procedure required by the Fifth Amendment before the measures were imposed.

The government can appeal, and the Pentagon still has every right to choose a different supplier.

The verdict therefore does not end the dispute; it only imposes clearer rules for the next round, in which it will be decided under what conditions the US military may use the most powerful AI systems.

The real problem is broader than a single contract with one company. The US military wants full operational control over the technology it deploys in intelligence, cyber operations, planning and other sensitive tasks.

Anthropic believes there are types of use for which today's models are not yet reliable enough, or that exceed the limits the company wants to uphold.

With ordinary software, such a dispute is resolved by changing the supplier. With a system that can become part of decision-making in war, the consequences are far more serious.

From partnership to open dispute

By the time the conflict began, Anthropic was already deeply involved in the US defence system. In July 2025, the Department of Defense awarded the company a two-year

development **contract**, capped at \$200 million, to apply advanced AI systems in national security.

Claude was introduced into classified environments, and the company listed intelligence analysis, modelling and simulation, operational planning, and cyber operations as areas of collaboration.

For **Anthropic**, the Pentagon was one of its most important potential customers, while for the department, Claude was one of the few most advanced models capable of operating in highly demanding environments.

The split occurred during negotiations over future terms of use. The Pentagon demanded the ability to use the model for all lawful purposes. Anthropic accepted a very wide range of military applications but retained two prohibitions.

The company did not want Claude to be used for mass surveillance of Americans or to be involved in fully autonomous systems that would apply lethal force without human decision-making.

The Pentagon viewed the problem through the lens of command responsibility

Anthropic argued that today's models are still susceptible to errors that can be corrected in ordinary commercial applications, whereas the same type of error in a target-selection or attack system could have irreversible consequences.

The Pentagon viewed the problem through the lens of command responsibility. If an operation is lawful and approved through the US chain of command, the department does not want the supplier subsequently to prohibit a particular use of the technology.

From a military planning perspective, such a request follows a clear logic. The armed forces cannot develop critical capabilities amid constant uncertainty over whether a supplier

will change the rules of use, restrict a model's functions, or decide that an operation no longer complies with its policy.

The legal problem arose only when the dispute over those conditions shifted from procurement and contracts into the realm of national security.

Where the Pentagon lost its legal footing

On 27 February, **President Trump** ordered federal agencies to stop using Anthropic's technology, with a transition period for systems in which it was already in use. That same day, Defence Secretary **Pete Hegseth** began the process of designating Anthropic as a supply chain risk.

Such a designation has consequences that extend far beyond the decision not to enter into further contracts with a company. It is part of a legal regime designed to protect military systems from compromise, sabotage and hostile influence.

In **Anthropic's case**, there was no allegation that the company had granted an adversary access to US systems, modified software for sabotage, or concealed a security flaw that would have endangered the armed forces.

Even after the dispute escalated, the department continued to rely on Claude for ongoing activities during the transition period

The gist of the dispute was public knowledge: Anthropic rejected two categories of use and accepted that the Pentagon might choose another supplier as a result. Judge Lin therefore concluded that the label "supply chain risk" was used outside the purpose for which it was provided in law and that the procedure was arbitrary and legally insufficient.

In doing so, the Pentagon made its own legal

position more difficult. Even after the dispute escalated, the department continued to rely on Claude for ongoing activities during the transition period, while simultaneously considering the continuation of certain forms of cooperation.

That hardly fits with the claim that the supplier poses the kind of danger for which the supply chain risk mechanism was introduced. The ruling leaves the Pentagon considerable freedom in choosing suppliers but does not allow it to reclassify a political or contractual dispute as a security threat without a proper basis.

Anthropic's legal victory does not resolve the main issue

Judge Rita Lin's ruling settled the question of the legality of the Pentagon's move, but not the far more important dispute over who determines the limits of the military's use of artificial intelligence.

The court did not order the Pentagon to use Claude or decide which rules are best for autonomous weapons, surveillance, or other military applications of artificial intelligence. The state's right to terminate cooperation with a supplier whose conditions do not suit it remains intact.

Civilian control of the armed forces is exercised by institutions answerable to voters, Congress and the courts

Separate procedures linked to other grounds for limiting cooperation also mean that the relationship between the two parties will not simply return to its **pre-conflict state** because one decision is overturned.

For the American state, the issue of operational sovereignty remains very serious. Civilian control of the armed forces is exercised by institutions answerable to voters, Congress and the courts.

Decisions about the use of force cannot permanently depend on how the management of a private company interprets its own ethical rules.

If a country concludes that, for a certain legitimate military capability, it needs a model without the limitations imposed by the supplier, the solution is to choose another product, develop its own system, or predetermine the rules by contract. Once this decision has been made, the next phase of the dispute will be based on it.

OpenAI has already shown that there is a different model

In February, **OpenAI** reached an agreement with the Pentagon allowing broad use of its systems for legitimate purposes, with safeguards relating to domestic surveillance and autonomous weapons.

That arrangement differs from what Anthropic sought, but it shows that the defence department can accept restrictions when they are precisely defined and when it is clear in advance who will decide how they are interpreted and applied.

Where is the line between permissible surveillance of a specific person and the mass collection of data?

Therefore, the claim that any limitation set by an AI company is automatically incompatible with military needs is no longer convincing. A much more difficult dispute will concern the substance of those restrictions.

Many important questions remain. Does the ban apply only when the AI independently decides to use lethal force, or also when the system suggests targets and the human merely formally approves its proposal?

Where is the line between permissible surveillance of a specific person and the mass

collection of data? And who determines that limit?

Such rules must be clearly defined before AI becomes an integral part of military operations, because by then it will be too late to resolve fundamental issues in the midst of a political or security disputes.

Military AI is changing the relationship between the state and its suppliers

A traditional defence contractor produces an aircraft, radar or missile system to agreed specifications. With the most advanced AI models, however, the relationship no longer ends at delivery.

Models are continuously updated, their capabilities change, they connect to new databases and other systems, and their performance depends on computing infrastructure that is often controlled by the private sector. The supplier can therefore remain technically relevant long after the product has entered use.

If a small number of private companies own models that the military cannot quickly replace or independently reproduce, their bargaining power becomes substantially greater than that of traditional software producers.

Contracts will need to describe specific situations, authorities and responsibilities in detail

With that power comes responsibility. When a model is used in an intelligence assessment, a cyber operation, or a decision related to the use of force, the company can hardly leave all the consequences entirely to the user.

The same kind of dispute will soon arise beyond surveillance and autonomous weapons. Biological research, offensive cyber operations, analysis of large sets of personal

data, selection of military targets, and systems that propose decisions to commanders in a short time will all require far more precise rules.

General statements about lawful use and responsible artificial intelligence will not suffice. Contracts will need to describe specific situations, authorities and responsibilities in detail.

The boundaries of military AI will have to be negotiated in advance

In the short term, the administration is likely to appeal, while the Pentagon will continue to reduce its reliance on Anthropic wherever it believes the company's restrictions hinder future plans.

The court cannot compel the military to choose Claude, and the political relationship between Hegseth and Anthropic's management is unlikely to normalise quickly. The company obtained protection from a type of state measure, but it did not secure the continuation of the contract or a return to its previous position at the Pentagon.



The court cannot compel the military to choose Claude, and the political relationship between Hegseth and Anthropic's management is unlikely to normalise quickly

The biggest consequence of this dispute will probably not be the verdict itself, but a change in how the Pentagon contracts for the use of the most advanced AI systems in future.

It will no longer be sufficient for the contract

to define only the price, access to technology and technical support.

It will also need to specify in advance who bears responsibility if the model fails, when a human must take over, what happens if the company changes the rules of use, and whether the military can continue the operation if the supplier withdraws support. This will be one of the key issues in future military planning and readiness.

The most realistic outcome is a compromise that will be less dramatic than the current political conflict but more consequential for the future.

The Pentagon will insist that decisions on military operations remain within the country's chain of command. AI companies will seek to ensure that their models are not used outside pre-contracted categories, particularly where legal liability is high and mistakes could be deadly.

The line between those two requirements will be defined by contracts, procurement standards and, increasingly, legislation that Congress must pass before technology practices advance further.

Judge Rita Lin's ruling did not determine who will ultimately have the final say when AI becomes an integral part of military decision-making. It only determined that the Pentagon could not resolve the issue. Declaring a recalcitrant supplier a security risk requires a solid legal basis, which the court did not find in this case.

The next serious conflict may arise when an AI system is already deeply embedded in an operation from which it cannot easily be withdrawn.

Until then, the US will have to decide how much technological dependence on private companies it can accept, and what conditions it must establish before it begins to rely on their models in a crisis or war.