

Una jueza veta la decisión



de Trump de enlistar a
Anthropic como "un riesgo
para la cadena de
suministro"

La jueza ha calificado como "ilegal e infundada" la decisión de considerar a Anthropic un riesgo para la cadena de suministro en materia de seguridad nacional.

El director ejecutivo de Anthropic, Dario Amodei, en la AI Impact Summit de Nueva Delhi. FOTO: LUDOVIC MARIN/GETTY IMAGES

Una jueza federal prohibió este jueves a la Administración Trump calificar a Anthropic como un riesgo para la seguridad nacional a raíz de una disputa sobre el uso de sus modelos de inteligencia artificial. La jueza dictaminó que dicha calificación constituía una **represalia inconstitucional**.

Medidas ilegales e infundadas

La resolución de la jueza federal de distrito Rita Lin, en California, anuló la decisión del 27 de febrero del secretario de Defensa, Pete Hegseth, de calificar a Anthropic como un **"riesgo para la cadena de suministro"**, además de inhabilitar a la empresa para optar a contratos federales. La sentencia también anuló una medida punitiva adicional impuesta por Hegseth que impedía a los contratistas o proveedores del ejército estadounidense hacer negocios con Anthropic. Lin llamó a la postura "arbitraria, caprichosa, un abuso de discrecionalidad y, por lo demás, contraria a la ley".

Nuestras historias más importantes, diario en tu correo electrónico.

Inscribirse

Ingresa tu dirección de correo electrónico para suscribirte a nuestro newsletter y formar parte del universo Wired. Tus datos personales serán tratados de acuerdo con nuestra [Convenio de Usuario](#) y [Aviso de Privacidad](#).

"Aunque es indiscutible que el Departamento de Guerra es libre de elegir al proveedor de IA que desee, las pruebas demuestran que las amplias medidas impuestas a Anthropic eran ilegales e infundadas", escribió Lin en un fallo de 59 páginas.

La jueza también dictaminó que nueve organismos (entre ellos el Pentágono, el Departamento del Tesoro, el Departamento de

Estado y el Departamento de Seguridad Nacional) habían impuesto sanciones indebidamente a Anthropic. La sentencia anula tales sanciones.

Sin embargo, Lin aseguró que el Pentágono no está obligado a usar los modelos de Anthropic, y que la sentencia no impide que el organismo opte por utilizar otros modelos. No fue posible contactar de inmediato con un portavoz del Pentágono para recabar comentarios, pero se espera que el departamento presente un recurso de apelación.

En un comunicado, la portavoz de Anthropic, Danielle Cohen, declaró: “Acogemos con satisfacción la sentencia del tribunal que declara ilegal esta designación de riesgo para la cadena de suministro. Seguimos centrados en colaborar de forma productiva con el Gobierno para aprovechar la IA en beneficio de nuestra seguridad nacional”.

Claude y sus aplicaciones militares

El caso surge de una amarga disputa surgida a principios de este año entre el Pentágono y Anthropic por un acuerdo de 200 millones de dólares para usar los modelos Claude del laboratorio de IA en aplicaciones militares.

La disputa comenzó tras los informes de que Estados Unidos había utilizado Claude en la operación para capturar al presidente venezolano Nicolás Maduro. Después de la operación, un empleado de Palantir transmitió a funcionarios estadounidenses las preocupaciones de un miembro del personal de Anthropic sobre cómo se habían usado sus modelos.

En las negociaciones, Anthropic insistió en establecer ciertos límites al uso de sus modelos, incluido el apoyo a armas autónomas letales y sistemas de vigilancia masiva. Pero Hegseth rechazó cualquier restricción, alegando que un contratista no podía dictar cómo se utilizaría la tecnología una vez implementada, insistiendo en que el contrato permitía “todo uso lícito”.

Tras el fracaso de las negociaciones en febrero, se consideró de forma generalizada que Hegseth estaba castigando a Anthropic al designarla como “riesgo para la cadena de suministro”. La medida de seguridad nacional incluía a la empresa en una lista negra que le impedía hacer negocios con el gobierno federal.

El Pentágono afirmó en aquel momento que la calificación de “riesgo para la cadena de suministro” era adecuada, ya que dar acceso a Anthropic a cualquier sistema clasificado “supondría un riesgo inaceptable” si el laboratorio de IA pudiera desactivar o alterar su tecnología, por ejemplo, en tiempo de guerra.

Contrademandas

Anthropic presentó dos demandas en respuesta, una ante el tribunal federal de distrito de California y otra ante el Tribunal de Apelación de los Estados Unidos para el Distrito de Columbia. En ellas que acusaba al Pentágono de violar las protecciones que le otorgan la Primera y la Quinta Enmienda por motivos ideológicos. El caso del Distrito de Columbia sigue en curso.

Los principales modelos de IA de Anthropic están considerados entre los mejores del mundo. Al margen de la designación, los modelos de la empresa han pasado recientemente a estar sujetos

al marco de supervisión de la IA de la Administración Trump debido a sus potentes capacidades.

Lin citó la participación de la Administración en la revisión de los modelos avanzados de la compañía como parte de su fallo.

“Incluso ahora, el gobierno está debatiendo la colaboración con Anthropic en su nuevo modelo, Mythos, en una serie de contextos sensibles”, escribió. “Nada de eso concuerda con un temor genuino a que Anthropic sea un saboteador que viciara su *software* para perjudicar la seguridad nacional”.

Información adicional de Maxwell Zeff.

Artículo originalmente publicado e