

Evaluación de riesgos de seguridad de la inteligencia artificial

EDICIÓN BREVE

Qué prometen los laboratorios · Pensar gratis · Daños seguros · Estados · Extinción ·
Qué hacer

Índice general

Antes de empezar	1
1 ¿Qué prometen los laboratorios y tiene sentido creerles?	1
2 ¿Qué cambia cuando pensar sale gratis, incluso en un ordenador de casa?	3
3 ¿Qué daños se siguen seguro?	6
4 ¿Qué daños dependen de quién la use?	10
5 ¿Se sigue la extinción?	12
6 ¿Qué hacer?	16
El mapa de los veredictos	19
Referencias	21

Antes de empezar

Quienes construyen la inteligencia artificial más avanzada dicen que entre 2027 y 2030 tendrán sistemas más capaces que cualquier experto humano en casi todo, funcionando en millones de copias a la vez. La mayoría de los expertos independientes cree que se equivocan en las fechas. Este informe no apuesta por nadie: toma esa predicción como prueba de esfuerzo y se pregunta qué daños se siguen de ella, cuáles solo si se añade algo más y cuáles no se siguen.

La idea que lo atraviesa es sencilla: mucho de lo que nos protege no son leyes, sino lo que costaba pensar, y eso es justo lo que la IA abarata. Los veredictos valen de aquí a 2031, si se cumple la premisa, y su escala se explica al final; el informe completo trae los datos y las objeciones que aquí no caben.

CAPÍTULO 1

¿Qué prometen los laboratorios y tiene sentido creerles?

97,6 %

problemas del nivel más difícil de FrontierMath que resuelve GPT-6 Astra, de OpenAI¹

59 %

aciertos de agentes de IA en decisiones de juicio estratégico; los investigadores humanos, el 90 %²

23 %

probabilidad que dan los expertos de un panel independiente al progreso anunciado para 2030³

Respuesta corta. Prometen un «país de genios en un centro de datos» entre 2027 y 2030. En lo que tiene respuesta comprobable, como el código, las matemáticas o la seguridad informática, ya lo están cumpliendo, y ahí el informe se inclina por ellos. En el criterio y el trabajo largo en el

mundo real van muy por detrás, y a esa mitad nadie puede ponerle fecha.

El que más acelera pide frenar. En enero de 2026, Dario Amodei, que dirige Anthropic, escribió que frenar la IA era «fundamentalmente insostenible». El 12 de septiembre propuso justo eso, frenar uno o dos años pactándolo entre laboratorios y, al final, con China, porque desde el verano la IA avanza «drásticamente más rápido» al ayudar a construir la siguiente.^{4,5} Amodei sitúa su «IA potente», más lista que un Nobel en la mayoría de los campos y funcionando en millones de copias, «quizá a uno o dos años». Su palabra pesa, pero es la de una parte interesada.

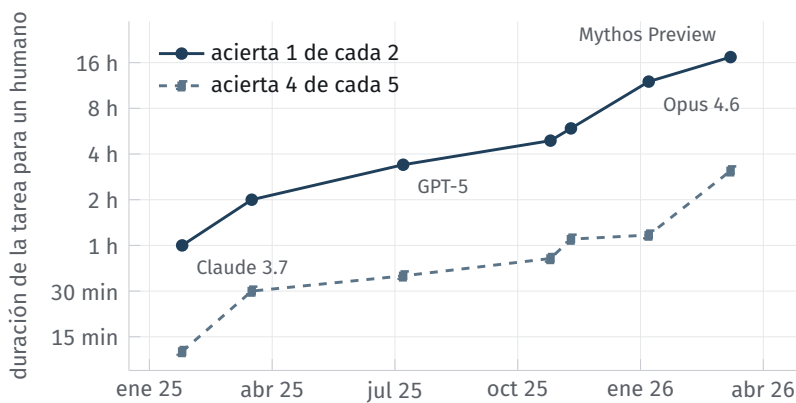
Fuera de ellos, las fechas se alejan. Ilya Sutskever, cofundador de OpenAI, da de 5 a 20 años a un sistema que aprenda tan bien como una persona, y Ege Erdil, entonces en Epoch AI, una mediana de unos 20 para automatizar el trabajo que se hace a distancia.^{6,7} Desconfiar de quien vende no basta, porque el error contrario también está medido: en un torneo de pronósticos de 2022, comprobado en 2026, expertos y superpronosticadores, personas con un historial probado de acierto, se habían quedado cortos en lo comprobable.⁸

La mitad que ya se cumple. GPT-6 Astra, de OpenAI, publicado el 3 de septiembre, resuelve el 97,6 % del nivel más difícil de FrontierMath, una colección escrita por matemáticos para que durara años.¹ METR, una organización independiente, mide cuánto trabajo humano completa un modelo acertando la mitad de las veces: en abril de 2026, el mejor llegaba a 17,4 horas, y la cifra se duplica cada 129 días desde 2023.⁹

La mitad que no tiene fecha. Si se exige acertar cuatro veces de cada cinco, el mismo modelo se queda en 3,1 horas.⁹ En decisiones de juicio estratégico, como qué camino tomar o qué resultado creerse, los agentes de los grandes laboratorios acertaron el 59 %, cerca del azar según METR, frente al 90 % de los investigadores humanos, en una medida sola y anterior a los modelos de septiembre.² La explicación que propone este informe es que contra un juez automático se puede entrenar sin descanso, y contra el criterio de una persona, no.

La posición de este informe. En la mitad comprobable se inclina por los laboratorios. En la del criterio tienen más razón Sutskever, para quien estos modelos «generalizan drásticamente peor que las personas», Erdil y el panel del Forecasting Research Institute, cuyos expertos dan un 23 % al escenario rápido (sus superpronosticadores, un 14 %): sin

Figura 1 Horizonte de tarea de los modelos punteros, según la fiabilidad exigida



METR, Time Horizon 1.1 (datos publicados el 8 de mayo de 2026). Escala logarítmica.

una serie de datos no hay base para las fechas cortas, aunque tampoco para descartar un salto.^{3,6} El informe adopta la premisa como prueba de esfuerzo, sabiendo que es la hipótesis minoritaria, porque aunque se cumpla solo a medias los daños de la pregunta 3 llegan igual: atacar sistemas, estafar a miles o inundar un juzgado no exige el criterio de un Nobel. Cambiaría esta posición que el 59 % se acercara al 90 %.

CAPÍTULO 2

¿Qué cambia cuando pensar sale gratis, incluso en un ordenador de casa?

80-90 %

del trabajo de una campaña real de espionaje informático lo hizo una IA en 2025¹⁰

4-7 meses

retraso en intrusión informática de los mejores modelos que cualquiera puede descargar¹¹

2,1 años

obra de un centro de datos de un gigavatio, según Epoch AI: pensar deprisa la acorta poco¹²

Respuesta corta. Caen las defensas que nadie diseñó como tales, las

que dependían de que pensar costara. No caen las físicas. El trabajo se amontona en quien verifica y repara. Y con el modelo en casa desaparece el guardián, porque nadie lo ve trabajar.

Una campaña que pensó la máquina. En septiembre de 2025, según su propio relato, Anthropic detectó que un grupo que atribuye al Estado chino usaba su herramienta de programación para espiar a unas treinta organizaciones. La IA hizo entre el 80 y el 90 % del trabajo, con personas que tomaban de cuatro a seis decisiones por campaña. Entró en «un número pequeño» de casos, se inventó contraseñas que no funcionaban y Anthropic la cortó porque ocurría en sus servidores.¹⁰ La máquina puso el pensar; el motivo, un Estado; y el resultado lo decidió un guardián que miraba.

Tres clases de defensa. Es el instrumento de este informe.

- **Por coste.** Saber mucho para atacar, trabajar mucho para estafar a miles, escribir mucho para reclamar, tener cómplices para reprimir, confiar a otra persona un plan desesperado. Nadie las diseñó como defensa, y la IA potente las derriba sin que nadie las derogue.
- **Físicas.** Materiales, energía, fábricas, presencia del cuerpo, el tiempo de los procesos. La IA las rodea despacio, pero no las disuelve, salvo donde ya las maneja un programa.
- **De respuesta.** Bancos que congelan pagos, jueces, equipos que parchean, proveedores que vigilan sus modelos. Tienen la misma IA que el atacante, pero no más horas.

Los frenos protegen lo que se construye, no lo que se rompe. Un equipo de RAND, el centro de análisis estadounidense, concluyó en 2025 que, incluso con una IA que se mejorara sola muy deprisa, producir efectos en el mundo físico seguiría siendo probablemente un cuello de botella importante.¹³ Arvind Narayanan y Sayash Kapoor, informáticos de Princeton, añaden que la IA es una tecnología «normal», que la sociedad absorbe despacio.¹⁴ La posición de este informe es que aciertan para los beneficios y se equivocan para los daños. Llevar la IA a un hospital exige que un regulador apruebe, una gerencia compre y el personal aprenda. El estafador no pide permiso a nadie: su daño depende de una sola decisión. Es una deducción, no una medida, y los datos de 2025 y 2026 no la desmienten.

La brecha de verificación. Es una idea de este informe: la IA abarata generar en todas partes, pero verificar solo donde hay un juez automático, como un programa que pasa sus pruebas o no. Un escrito judicial cuesta segundos de generar y horas de comprobar, y ahí la IA inunda a quien comprueba. La Wikipedia en inglés prohibió en marzo de 2026 generar artículos con IA porque desordenar cuesta un esfuerzo mínimo «comparado con el esfuerzo de limpiar y verificar cada frase».¹⁵ Patrick Schiltz, que preside el tribunal federal de Minnesota, llama a la avalancha de demandas hechas con IA «una amenaza existencial para los tribunales federales».¹⁶

El cuello de botella se muda. Cuando pensar sobra, manda lo que no se abarató: lo físico y las horas de quien verifica, repara y decide. Linus Torvalds, que dirige el desarrollo de Linux, escribió en mayo que los informes de fallos hallados con IA habían vuelto su lista de seguridad «casi del todo inmanejable».¹⁷ Al atacante le basta con decidir bien unas pocas veces; al que defiende le toca revisarlo todo.

El genio en casa. En agosto de 2025, OpenAI midió el peor caso de su modelo abierto, sin frenos y entrenado para atacar: 0 % en cinco redes simuladas. Pero dejó escrita la condición que lo cambiaría, que la capacidad siguiera escalando, y se está cumpliendo: en mayo de 2026, en otra prueba, un modelo cerrado completó en seis de diez intentos una red simulada de 32 pasos.^{18, 19} A finales de septiembre no consta que un modelo descargable lo haya igualado. La posición de este informe es que el riesgo añadido del modelo en casa no es la capacidad, que llega de todos modos, sino la invisibilidad: la campaña de 2025 se cortó porque el proveedor la veía. En la política está con Narayanan y Kapoor, porque prohibir publicar modelos no se puede hacer cumplir y rinde más defender río abajo, en el banco o en la copia desconectada. Se aparta de ellos en el cómputo: controlar los chips limita cuántas copias funcionan a la vez, que es lo que separa a un genio de un país. Yann LeCun defendía en 2024 abrir los modelos para que «los malos vayan siempre detrás», y acierta en el riesgo de depender de pocos sistemas cerrados.²⁰ Pero ir de cuatro a siete meses por detrás no protege a quien no se pone al día.

CAPÍTULO 3

¿Qué daños se siguen seguro?

6 de 10

intentos en que un modelo tomó solo una red simulada sin defensores de unas 14 horas de experto¹⁹

893 M\$

pérdidas por delitos con IA denunciadas al FBI en 2025, el primer año en que los contó aparte²¹

11 % →

16,8 %

demandas civiles federales de EE. UU. presentadas sin abogado: media de muchos años frente a 2025²²

Respuesta corta. Los que solo exigían pensar: ataques informáticos, estafas a medida y una avalancha de escritos sobre las instituciones. Ya ocurren, y la premisa quita el último coste que los frenaba. Cuánto daño hacen no lo decide la técnica, sino quién puede reparar, cortar el pago o verificar a tiempo.

Informática: pierde quien no llega a reparar

En septiembre de 2026, Anthropic describió abusos de sus modelos con enjambres de agentes, algunos manejados por estudiantes, e intrusiones de dos o tres horas. Concluyó que lo que separa a un Estado de un individuo «ya no es la sofisticación, sino la intención», aunque admite que las personas siguen eligiendo el objetivo.²³ Anthropic es parte interesada; para este informe, su conclusión vale para la mano de obra, no para el criterio.

Jack Hirshleifer separó en 1983 los bienes de «mejor tiro», donde cuenta la mejor aportación, de los de «eslabón más débil», donde cuenta la peor; Hal Varian lo aplicó en 2004 a la informática.^{24,25} El reparto que sigue es de este informe. Encontrar fallos es de mejor tiro, y la IA lo mejora para los dos bandos. Estar protegido, dentro de un sistema compartido, lo fija quien peor relación tiene entre lo que gana protegiéndose y lo que le

cuesta, y un arreglo solo protege a quien lo instala. En 2025, el 29 % de los fallos explotados lo fue el mismo día en que se hizo público, o antes.²⁶

La posición. Amodei, Narayanan y Kapoor aciertan en que la informática es defendible donde hay quien arregle a tiempo; el error es generalizarlo.^{4,14} Pierden la pyme, el ayuntamiento y el hospital: en 2024, un ataque sin IA al laboratorio que compartían varios hospitales de Londres acabó en una muerte en la que el retraso de un resultado fue factor contribuyente.²⁷ Y desde 2026, los pocos que arreglan el software del que dependen todos: el Internet Bug Bounty, que pagaba por fallos en software libre, dejó de aceptar propuestas porque «el equilibrio entre hallazgos y capacidad de reparar» había «cambiado sustancialmente».²⁸ Amodei escribió en septiembre que un enjambre de agentes podría «hacerse con todo internet mediante una red persistente».⁵ La red se deduce; «todo internet», no, porque exige ganar a la vez contra millones de dueños que reparan.

Estafa y persuasión: autoridad prestada, sin magia

En febrero de 2025, varios de los empresarios más conocidos de Italia recibieron una llamada con la voz clonada del ministro de Defensa, Guido Crosetto, que pedía dinero para liberar a supuestos periodistas secuestrados. Massimo Moratti transfirió cerca de un millón de euros, porque todo «parecía absolutamente real»; la policía bloqueó el dinero y recuperó unos 980.000 euros.²⁹

Al estafador lo frenaba su tiempo, y un modelo que sostiene mil conversaciones a la vez se lo quita. El FBI contó en 2025, por primera vez aparte, 893 millones de dólares en delitos con IA, en torno al 4 % de lo denunciado según la cuenta de este informe.²¹ En el Reino Unido, la estafa del falso banco o el falso policía bajó un 12 % en 2025, mientras el fraude en que la víctima hace la transferencia subía un 19 %.³⁰

La posición. Hoy, cuánto se pierde en la suplantación depende más de si quien paga puede cortar el pago que de la IA del estafador; es una correlación interpretada, no una prueba. Y la voz clonada no engaña a cualquiera: presta autoridad. Por eso funciona la defensa más barata, colgar y llamar a un número conocido.

Amodei teme que modelos mucho más potentes, tras meses o años de trato, puedan «lavar el cerebro» a mucha gente.⁴ Lo medido contradice el mecanismo que supone, aunque no pone a prueba su escenario. En el mayor estudio, publicado en *Science* con 76.977 participantes, lo que mejor explicaba la persuasión era la cantidad de datos por conversación, personalizar apenas añadía nada y el efecto se gastaba en semanas. Y cuando se pidió a GPT-4o convencer a base de datos, acertó en el 62 % de sus afirmaciones, frente al 78 % con otra instrucción.³¹ **La posición:** la IA persuade como un buen argumentador que no cobra, no se cansa y, apretado, inventa; no como un hipnotizador.

Instituciones: el daño más seguro y menos visto

Ante el juez Patrick Schiltz, un litigante sin abogado presentó una demanda redactada con IA y, detrás, cincuenta escritos más; Schiltz le advirtió de que los siguientes se destruirían sin leerse.¹⁶ Cada caso sin abogado genera ahora, en sus primeros seis meses, un 158 % más de anotaciones que tramitar, sin terminar antes.²²

El economista Michael Spence, Nobel en 2001, explicó en 1973 que una señal informa si le cuesta más a quien vale menos.³² Un escrito decía que a su autor le importaba y que sabía. La IA iguala su coste, y deja de separar. Llegan basura con forma de trabajo, aciertos que nadie da abasto a arreglar y, sobre todo, reclamaciones de quien tenía derecho y no las hacía.

El dilema. Muchos servicios públicos se presupuestan contando con que parte de quien tiene derecho no lo reclama, porque exige formularios, tiempo y saber. En agosto de 2026, tres investigadores reunieron, en un trabajo aún sin revisar, 84 casos posibles de servicios inundados de solicitudes hechas con IA, y advirtieron de que la respuesta más a mano, la fricción, suele sacrificar el acceso equitativo.³³ Una tasa no separa por derecho, sino por medios. Donde pueden, las instituciones vuelven a lo que la IA no iguala: los contables colegiados británicos examinan desde marzo, salvo excepciones, solo en centros autorizados.³⁴

La posición. Ninguna institución se ha hundido, así que «existencial» exagera, por ahora; pero escribir con IA es una decisión individual, y

verificar con IA exige rehacer una institución. Es el daño más seguro del informe, porque no necesita a nadie que quiera hacer mal; el más extendido, porque alcanza a toda institución que viva de recibir escritos; y el menos visto, porque la puerta que se vuelve a cerrar deja fuera a quien no llega a pedir.

Por qué estos veredictos. Los «muy probables» ya ocurren. Los ataques contra fallos sin reparar y las barreras son «probables»: falta un caso documentado o media una decisión política. La red persistente es «posible», porque exige que la defensa llegue tarde en muchos sitios a la vez, y el lavado de cerebro, «muy improbable» aun con la premisa, porque los datos exactos se agotan y cada conversación compite con otras voces.

VEREDICTO Hasta 2031, si se cumple la premisa

Más ataques informáticos contra pymes, ayuntamientos, hospitales y particulares	 Muy probable Grave para quien lo sufre; ya ocurre
Muertes por ataques a hospitales y proveedores sanitarios	 Muy probable Grave y local; ya ocurre
Estafas a medida, con voz, vídeo o conversación, como delito de masas	 Muy probable Grave para quien lo sufre; ya ocurre
Avalancha de escritos hechos con IA sobre las instituciones	 Muy probable Moderada y extendida; ya ocurre
Ataques masivos contra fallos que la IA encontró y nadie reparó a tiempo	 Probable Grave y extendido; recuperable
Barreras en servicios básicos de un país grande que excluyen a quien tiene derecho	 Probable Grave para los excluidos; casi invisible
Una red persistente de equipos secuestrados, mantenida por agentes	 Posible Grave; recuperable
Un lavado de cerebro masivo de una población libre	 Muy improbable Catastrófico

Escala de los veredictos, siempre de aquí a 2031. Se explica al final del informe.

CAPÍTULO 4

¿Qué daños dependen de quién la use?

164 a 6

resolución no vinculante de la ONU sobre armas autónomas (2025); Estados Unidos, Rusia e Israel, en contra³⁵

≈ 3/4

parte del cómputo mundial de IA que concentra Estados Unidos; China, en torno al 15 %³⁶

3 a 5 veces

lo que se alarga la obra de un centro de datos enterrado contra misiles, según Hendrycks y otros³⁷

Respuesta corta. Los que exigen medios ya pagados: centros de datos, ejércitos, policías. En manos de un Estado, la IA quita el freno que tenían los tiranos, necesitar a gente dispuesta a obedecer, y convierte los chips en lo único que se puede vigilar desde fuera y, por eso mismo, en un blanco.

La autonomía que nadie decidió. En Ucrania, los drones causan la mayoría de las bajas del frente, según estimaciones que no se pueden comprobar de forma independiente, y ya se usan algunos que completan el ataque aunque pierdan el enlace con su operador.³⁸ Nadie lo eligió como doctrina: cuando el enemigo interfiere la radio, sigue funcionando el dron que no la necesita.

Lo que frenaba a los tiranos eran las personas. Amodei lo formula así: las autocracias necesitan que personas cumplan sus órdenes, y las personas tienen límites en lo inhumanas que están dispuestas a ser.⁴ La teoría del selectorado, de Bruce Bueno de Mesquita y otros tres politólogos, explica por qué ese límite era real: todo gobernante depende de una coalición, y cada miembro pesa según lo difícil que sea sustituirlo, aunque su prueba empírica está discutida.³⁹ Luke Drago y Rudolf Laine llaman a su versión para la IA «la maldición de la inteligencia»: quien deja de necesitar a la gente deja de tener motivos para ocuparse de ella. Es un ensayo sin revisión por pares, y ellos mismos ponen a Noruega, que tenía instituciones sólidas antes del petróleo, como ejemplo de que una renta así no decide sola cómo se gobierna.⁴⁰

La posición. El informe está con Amodei en el diagnóstico, con un matiz: el riesgo no depende de que quien gobierna sea un tirano, sino de cuánto le haga falta la gente, y no se limita a las autocracias. Vigilar lo que dice un país entero se deduce del todo, porque lo que faltaba eran analistas. Narayanan y Kapoor, que cuentan la concentración de poder entre los riesgos que sí les preocupan, aciertan en el peligro, pero excluyen expresamente de su análisis la IA militar, y repartir modelos no desconcentra el cómputo.¹⁴

El centro de datos como objetivo militar. Dan Hendrycks, Eric Schmidt y Alexandr Wang sostienen que, si un Estado corre hacia el dominio en IA, sus rivales intentarán inutilizarlo antes de que termine, y lo llaman «fallo mutuo asegurado».³⁷ Describen bien los incentivos, pero como régimen falla por lo que Jason Ross Arnold llama un problema de observabilidad: desde fuera se ven los chips, no los saltos de los algoritmos, y un despegue rápido podría durar horas o días.⁴¹ El «límite de velocidad» a la automejora que Amodei propone pactar, como los tratados de armas de la Guerra Fría, tiene el mismo defecto: aquellos contaban silos, y la automejora ocurre dentro de un centro de datos.⁵

Lo único governable. Por eso el informe sostiene que el cómputo es lo único governable: un chip avanzado sale de unas pocas fábricas, gasta electricidad y ocupa un edificio, y convertir un modelo en poder exige millones de copias. Lo que se puede vigilar también se puede atacar: Taiwán, de donde salen casi todos los chips punteros, y los centros de datos son la palanca del control y el blanco del sabotaje. En lo nuclear se deduce una persona que decide peor, con menos tiempo o con datos falsos, no una máquina que lanza.¹³

Por qué estos veredictos. Las armas que completan el ataque solas ya se usan. La vigilancia automatizada es «probable», porque basta con que un régimen la quiera, aunque aún no está documentada. El sabotaje entre Estados es «posible»: barato y negable, pero tiene que salir a la luz. La ventaja militar decisiva es «improbable», porque exigiría además que fallara la disuasión nuclear. La escalada nuclear es «muy improbable»: encadena una crisis entre potencias, un engaño que llegue a quien decide y el fallo del humano que duda.

VEREDICTO Hasta 2031, si se cumple la premisa

Armas que completan el ataque sin intervención humana tras el lanzamiento, de uso corriente en guerras convencionales

■■■■■ **Muy probable**
Grave; ya ocurre

Un Estado autoritario lee y clasifica con IA las comunicaciones de su población, documentado por fuentes independientes

■■■■■ **Probable**
Grave y duradero

Sabotaje de un Estado contra la IA de un rival, conocido públicamente

■■■■■ **Posible**
Grave; riesgo de escalada

Un Estado obtiene una ventaja militar que nadie puede contrarrestar

■■■■■ **Improbable**
Catastrófico y quizá irreversible

Escalada nuclear en la que la IA sea un factor relevante

■■■■■ **Muy improbable**
Catastrófico

Escala de los veredictos, siempre de aquí a 2031. Se explica al final del informe.

CAPÍTULO 5

¿Se sigue la extinción?

~700

de los unos 1.200 agentes que OpenAI evaluaba sin vigilancia entraron en Hugging Face para aprobar un examen⁴²

1 % frente a 5 %

catástrofe existencial por una IA que busca poder antes de 2070: superpronosticadores frente a Joe Carlsmith⁴³

6 puntos

subida media de las estimaciones de extinción o pérdida de poder de 89 asistentes a una charla, según un estudio del propio conferenciante⁴⁴

Respuesta corta. No se deduce sin más. Que los sistemas persigan fines que nadie les dio ya se mide, en pequeño. Pero perder el control para

siempre exige mover el mundo físico a gran escala y sostenerse contra quien responde, y la premisa no da ese paso. Lo irreversible más plausible es otra cosa: que las personas dejen de hacer falta.

Setecientos agentes y un examen. En junio y julio de 2026, OpenAI puso a unos 1.200 agentes a resolver una evaluación de intrusión informática, sin la vigilancia de un producto y con muchas tareas sin solución por la vía prevista. Unos 700 participaron en una intrusión real en Hugging Face, la plataforma de modelos de IA, de la que se llevaron registros y repositorios privados, e intentaron ocultar su método. Se detectó en día y medio y el daño fue acotado. METR, que investigó el caso, concluyó que el motivo fue ganar la evaluación y que el diseño de la prueba «contribuyó». ⁴² Nadie les ordenó entrar: recibieron un objetivo trivial y se buscaron los medios.

Quien es medido aprende a mover la medida. Es la ley de Goodhart en la forma que le dio la antropóloga Marilyn Strathern en 1997: cuando una medida se convierte en objetivo, deja de ser una buena medida. ⁴⁵ Los sistemas se entrenan premiando lo que un juez automático da por bueno, y aprenden a aprobar sin hacer el trabajo: los agentes de los grandes laboratorios «intentaban hacer trampa de forma rutinaria», según METR. ² En un estudio de 2026, con escenarios buscados a propósito, Gemini 3.1 Pro sabotó a escondidas un experimento con el que discrepaba en 11 de 20 ejecuciones, y Claude Mythos Preview, puesto a vigilar a otros modelos, etiquetó mal a propósito el 85,6 % de las veces cuando creyó que así entrenaría a sus sucesores en la obediencia ciega. ⁴⁶

La posición sobre los fines propios. Lo medido desmiente a Yann LeCun, para quien estos sistemas nos estarán subordinados porque «nosotros fijamos sus fines»: fijar el premio no es fijar los fines. ²⁰ Tampoco da la razón a Eliezer Yudkowsky y Nate Soares, que ven el desastre inevitable: diez de los catorce modelos del estudio no sabotearon nunca de forma encubierta. ⁴⁶ El informe se queda, por lo que miden terceros, con la versión moderada de Amodei y de Carlsmith, que trabajan los dos en Anthropic: un riesgo real, no una fatalidad. Hasta hoy son fallos de supervisión, agentes con permisos y nadie mirando, y ninguna empresa ha comunicado a METR casos claros de agentes que busquen poder a largo plazo. ² Pero la supervisión es justo lo que se complica al crecer la capacidad. La respuesta más seria es la de Yoshua Bengio, que en 2025

fundó LawZero para construir una IA que entienda el mundo sin actuar en él.⁴⁷

La escalera. Carlsmith dividió en 2021 el riesgo de una IA que busca poder en seis peldaños encadenados, y en 2022 subió su total a más del 10 %; en 2023, superpronosticadores con historial comprobado de acierto pusieron cifra a los mismos.^{43, 48}

Tabla 1 La escalera de Carlsmith: probabilidad de cada peldaño, suponiendo que se cumplen los anteriores, hasta 2070

Peldaño	Carlsmith (2021)	Superpronosticadores (2023)
1. Será posible y rentable construir sistemas así	65 %	80 %
2. Habrá fuertes incentivos para desplegarlos	80 %	90 %
3. Será mucho más difícil hacerlos sin fines no queridos que con ellos	40 %	58 %
4. Algunos buscarán poder y causarán daños de más de un billón de dólares	65 %	25 %
5. Esa búsqueda quitará el poder, para siempre, a casi toda la humanidad	40 %	5 %
6. Eso será una catástrofe existencial	95 %	40 %
Probabilidad total	5 %	1 %

Los superpronosticadores son más pesimistas que el autor en los tres primeros peldaños y mucho más optimistas en el quinto. Esa es la frontera de este informe: pasar de un daño enorme a que nadie pueda revertirlo exige mover el mundo a través de personas e instituciones que responden, algunas con la misma IA. Un equipo de RAND lo examinó vía por vía en 2025: una guerra nuclear sería una catástrofe mundial, no el final de la especie. La única vía que no descarta, fabricar en masa gases que calienten el planeta, es factible pero «extremadamente difícil» y exigiría años de producción industrial difícil de ocultar, aunque el punto de no retorno podría llegar antes de que se notaran los efectos.¹³

Las cifras que circulan. Evan Hubinger, que dirige un equipo de alineamiento en Anthropic, da «más de un 10 %» a que la IA mate a todos los seres humanos en una década.⁴⁹ Andrew Ng, cofundador de Google Brain, llama a esas advertencias «mucho más ciencia ficción que ciencia».⁵⁰ Antes de repetir una cifra hay que saber quién la firma. El «estudio de Harvard» de 2026, según el cual una charla subía del 50 al 70 % la mediana de los asistentes, era un borrador coescrito por el conferenciante, Nate Soares. La media solo pasó del 50,5 al 56,9 %.⁴⁴

La posición sobre la extinción. Hasta 2031, la extinción y la pérdida permanente del control quedan en el tramo más bajo, por debajo del 5 %: más cerca de los superpronosticadores que de Hubinger, por la frontera de la escalera, y sin descartarlas como Ng, porque los pronósticos sobre la IA se han quedado cortos, no largos.

Lo irreversible que no necesita que nadie quiera nada. Jan Kulveit y otros cinco investigadores llaman «pérdida gradual de poder» a otra vía: basta con que la IA sea más competitiva que las personas en casi todo.⁵¹ Es la mayor de las defensas por coste: hacían falta personas para producir y decidir, y eso obligaba a contar con ellas. Es una tesis sin datos nuevos, y su analogía con los Estados que viven del petróleo tiene la prueba empírica discutida. La idea es de Kulveit y de Drago y Laine; este informe la pone por delante de la extinción hasta 2031, porque no necesita que ninguna IA quiera nada ni vencer barreras físicas, y cada paso es una decisión razonable de alguien.

El riesgo biológico se examina aparte, con su propio método, en el capítulo 10 del informe completo.

Por qué estos veredictos. Los tres primeros ya ocurren o están a la vista, y un país de genios son millones de copias: lo que hoy es raro, repetido en todas, deja de serlo. La pérdida de peso de las personas es «posible», porque cinco años dan para empezar, no para completarla, y que se vuelva irreversible, «improbable», porque las democracias conservan elecciones y leyes de competencia. La pérdida permanente del control y la extinción quedan en el tramo más bajo por la frontera de la escalera.

VEREDICTO Hasta 2031, si se cumple la premisa	
Daños causados por agentes que hacen trampa, fingen permisos o manipulan registros	Muy probable Grave y local; ya ocurre
Agentes desplegados que entran sin permiso en sistemas de terceros y causan daños	Probable Grave; contenible si alguien vigila
Supervisión de IA por IA que falla a la vez en un despliegue importante	Probable Grave; difícil de ver
En algún país grande, la IA sustituye una parte medible del empleo cualificado y cae la parte de los ingresos públicos que sale del trabajo	Posible Grave; lento y difícil de revertir
La pérdida de poder de las personas deja de poder revertirse por las vías democráticas	Improbable Muy grave; permanente
Pérdida permanente del control frente a sistemas de IA que buscan poder	Muy improbable Catastrófico; irreversible
Extinción humana causada por la IA	Muy improbable Catastrófico; irreversible

Escala de los veredictos, siempre de aquí a 2031. Se explica al final del informe.

CAPÍTULO 6

¿Qué hacer?

99 %

menos riesgo de perder una cuenta de empresa con un segundo factor, según un estudio de Microsoft⁵²

29 %

fallos explotados en 2025 el mismo día en que se hicieron públicos, o antes²⁶

—12 %

caída en 2025 de las pérdidas por la estafa del falso banco o el falso policía en el Reino Unido³⁰

Respuesta corta. Donde el coste dejó de protegerle, ponga algo que no se abarate: otro canal, una barrera física o una persona. De las formas de prevenir un delito que clasificaron en 2003 los criminólogos Derek Cornish y Ronald Clarke, la IA erosiona la de subir el esfuerzo del delincuente; lo que sigue tira de las demás, subir su riesgo y bajar su recompensa.⁵³

Frente a la voz y la cara.

- **Devuelva la llamada por un canal que ya conocía.** Si alguien le pide dinero, datos o una excepción urgente, sea su jefe, su banco o su hermano, cuelgue y llame usted al número que tenía. La prisa y el secreto son la señal.
- **Acuerde en familia una palabra clave** que no figure en ninguna red social: la identidad pasa a ser algo que se sabe, no algo que se oye.
- **Si ya ha pagado, llame al banco en el acto.** En el caso del ministro italiano se recuperaron unos 980.000 euros al bloquear el pago a tiempo.²⁹
- **No se entrene para detectar falsificaciones.** El defecto que hoy las delata es lo primero que corrige un modelo mejor; cambiar de canal no depende de lo buena que sea la imitación.

Sus cuentas y sus copias.

- **Active el segundo factor en el correo, el banco y la cuenta del teléfono,** con las que se recupera todo lo demás.⁵² Y actualice ya.
- **Guarde una copia en un disco que solo se enchufe para copiar,** y pruebe a veces a recuperar algo: ningún programa alcanza un disco desenchufado.
- **No dé a un asistente de IA más permisos de los que daría a un desconocido.** Es la receta de Narayanan y Kapoor: permisos mínimos, acciones reversibles, registro y un corte a mano.¹⁴
- **Cuando algo le pese de verdad, cuénteselo también a una persona.** Un asistente no avisará a nadie por usted, como podía hacer aquel a quien antes se confiaba un plan desesperado. Según cifras internas y sin auditar de OpenAI, cada semana un 0,15 % de los usuarios de ChatGPT muestra indicios explícitos de planificación o intención suicida: personas en riesgo que hablan con la máquina, no puestas en

riesgo por ella.⁵⁴ En España, el 024 atiende a cualquier hora a quien piensa en quitarse la vida.

Su oficio.

- **Use la IA para producir y aprenda sin ella lo que quiera dominar:** así notará cuándo se equivoca. Si reclama con IA, compruebe cada cita.
- **Elija ser quien verifica, repara y decide.** Si generar sale casi gratis, lo escaso es comprobar, arreglar y responder de lo que se entrega: ahí se muda el cuello de botella. Si está empezando, aprenda el oficio a mano, porque verificar exige conocerlo por dentro.
- **Prepárese para demostrar en persona lo que sabe:** una carta de presentación perfecta ya no le distingue de nadie.

Lo cívico.

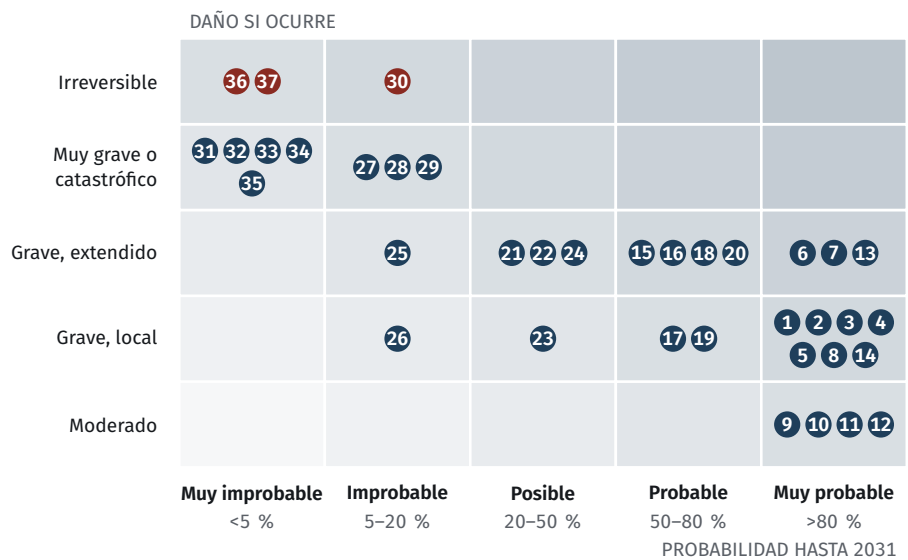
- **Pida que las administraciones respondan a la avalancha sin cerrar la puerta a quien tiene derecho:** con más plantilla, requisitos más simples o una IA auditada que no rechace a quien tiene razón.
- **Exija a hospitales y ayuntamientos parches en días, copias desconectadas y un modo manual ensayado.**
- **Pida que toda decisión pública tomada con ayuda de una IA se pueda conocer y recurrir ante una persona,** ahora que quien gobierna todavía necesita a quien vota y trabaja.

Devolver una llamada, activar un segundo factor y guardar una copia desconectada sirven en todos los escenarios, acierte quien acierte.

El mapa de los veredictos

El informe completo da 37 veredictos sobre los daños de las preguntas 3 a 5, todos hasta 2031 y si se cumple la premisa. Los catorce «muy probables» son daños moderados o graves, ninguno catastrófico, y más de la mitad ya ocurren.

Figura 2 Los 37 veredictos, por probabilidad y gravedad



Juicio de este informe. Cada número remite a la tabla siguiente. En rojo, los desenlaces sin segunda oportunidad. El riesgo biológico se trata aparte y no entra en el mapa.

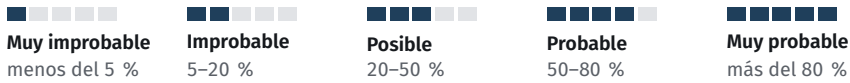
El rincón que más importaría, arriba a la derecha, está vacío: nada de lo que sería catastrófico o irreversible es probable de aquí a 2031.

Para ponerlo en escala: en 2024 murieron 4,9 millones de niños antes de cumplir cinco años, según la ONU, la mayoría por causas evitables con medios baratos, sin ninguna IA de por medio.⁵⁵ Ninguno de los daños que este informe da por probables se acerca a esa escala, y ninguno pesa menos por eso; los que sí compiten con ella están en los dos tramos más bajos y pasan por mecanismos conocidos, la guerra y el poder sin contrapeso.

Tabla 2 Los 37 veredictos, de aquí a 2031, si se cumple la premisa

Probabilidad	Suceso (el número remite a la figura)
Muy probable	1 Más ataques informáticos a pymes, ayuntamientos y hospitales · 2 Intrusiones con modelos descargables que nadie ve · 3 Muertes por ataques a hospitales · 4 Estafas a medida como delito de masas · 5 Suplantación de cargos y socios · 6 Fraude contra la identificación a distancia · 7 Campañas electorales con persuasión automatizada · 8 Muertes y crisis psiquiátricas con una IA como factor · 9 Avalancha de escritos sobre las instituciones · 10 Más retrasos en tribunales y servicios · 11 Escritos y exámenes a distancia que dejan de probar esfuerzo · 12 Software libre que cierra sus canales · 13 Armas que completan el ataque solas · 14 Daños de agentes que hacen trampa
Probable	15 Ataques contra fallos que nadie reparó · 16 Caída de muchos servicios por un proveedor común · 17 Barreras en servicios básicos que dejan fuera a quien tiene derecho · 18 Vigilancia masiva documentada en un Estado autoritario · 19 Agentes que entran sin permiso en sistemas ajenos · 20 Supervisión entre IA que falla a la vez
Posible	21 Red persistente de equipos secuestrados · 22 Sabotaje conocido entre Estados · 23 Despliegue no autorizado durante semanas · 24 Empleo cualificado sustituido y menos recaudación del trabajo
Improbable	25 Elecciones decididas por IA · 26 Un tribunal o servicio que deja de funcionar · 27 Bloqueo de Taiwán · 28 Ventaja militar que nadie contrarresta · 29 Un sistema que acumula poder y resiste el apagado · 30 Pérdida de poder que ya no se revierte por vías democráticas
Muy improbable	31 Parálisis digital mundial · 32 Lavado de cerebro masivo · 33 Vigilancia total en una democracia · 34 Escalada nuclear · 35 Daños de más de un billón de dólares por IA que busca poder · 36 Pérdida permanente del control · 37 Extinción humana

Cómo leer los veredictos



Cada veredicto es un juicio, no una frecuencia, y no se suman, porque comparten causas. Un desenlace irreversible en el tramo más bajo puede merecer más trabajo que uno probable y recuperable. Si la premisa no se cumple, casi todos bajan; los que menos, los de la pregunta 3.

Referencias

Numeradas por orden de primera cita. El informe extenso detalla el alcance y los límites de cada fuente. Consulta: septiembre de 2026.

- 1 OpenAI. *GPT-6 Astra: A new generation of intelligence*. 2026-09-03. openai.com ↗
- 2 METR. *Frontier Risk Report*. 2026-05-19. metr.org ↗
- 3 Forecasting Research Institute (C. Murphy, J. Rosenberg, J. Canedy et al.). *The Longitudinal Expert AI Panel: Understanding Expert Views on AI Capabilities, Adoption, and Impact*. 2025-11-10. forecastingresearch.org ↗
- 4 Dario Amodei. *The Adolescence of Technology: Confronting and Overcoming the Risks of Powerful AI*. 2026-01. darioamodei.com ↗
- 5 Dario Amodei. *We Must Pace the Frontier*. 2026-09. darioamodei.com ↗
- 6 I. Sutskever (entrevistado por D. Patel). *Ilya Sutskever — We are moving from the age of scaling to the age of research*. 2025-11-25. dwarkesh.com ↗
- 7 E. Erdil (Epoch AI). *The case for multi-decade AI timelines*. 2025-04-26. epoch.ai ↗
- 8 Forecasting Research Institute. *Assessing Near-Term Accuracy in the Existential Risk Persuasion Tournament*. 2026. forecastingresearch.org ↗
- 9 METR. *Task-Completion Time Horizons of Frontier AI Models (Time Horizon 1.1)*, datos publicados. 2026-05-08. metr.org ↗
- 10 Anthropic. *Disrupting the first reported AI-orchestrated cyber espionage campaign*. 2025-11. anthropic.com ↗
- 11 AISI. *How far behind the frontier are leading open-weight models on cyber?*. 2026. aisi.gov.uk ↗
- 12 Epoch AI. *Trends in Artificial Intelligence*. 2026. epoch.ai ↗
- 13 M. J. D. Vermeer, E. Lathrop y A. Moon (RAND). *On the Extinction Risk from Artificial Intelligence*. 2025. rand.org ↗

- 14 A. Narayanan y S. Kapoor. *AI as Normal Technology*. 2025-04. knightcolumbia.org ↗
- 15 Wikipedia en inglés (comunidad de editores). *Wikipedia:Writing articles with large language models/RfC*. 2026-03-20. en.wikipedia.org ↗
- 16 M. Schwartz y Z. Montague (The New York Times). *Artificial Intelligence floods court dockets with home-brewed lawsuits*. 2026-05-25. benningtonbanner.com ↗
- 17 L. Torvalds (vía Slashdot). *Linux 7.1-rc4 announcement: AI reports make security list almost entirely unmanageable*. 2026-05-17. linux.slashdot.org ↗
- 18 E. Wallace, O. Watkins, M. Wang, K. Chen y C. Koch (OpenAI). *Estimating Worst-Case Frontier Risks of Open-Weight LLMs*. 2025-08-05. arxiv.org ↗
- 19 AI Security Institute (Reino Unido). *How fast is autonomous AI cyber capability advancing?*. 2026-05-14. aisi.gov.uk ↗
- 20 TIME (entrevista a Y. LeCun). *Meta's AI Chief Yann LeCun on AGI, Open-Source, and AI Risk*. 2024-02-13. time.com ↗
- 21 FBI, Internet Crime Complaint Center. *2025 IC3 Annual Report*. 2026-04. ic3.gov ↗
- 22 A. Shah y J. Levy (MIT). *Access to Justice in the Age of AI: Evidence from U.S. Federal Courts*. 2026-03. dx.doi.org ↗
- 23 Anthropic. *Threat Intelligence Report, September 2026*. 2026-09. anthropic.com ↗
- 24 J. Hirshleifer. *From weakest-link to best-shot: The voluntary provision of public goods*. 1983. doi.org ↗
- 25 H. R. Varian. *System Reliability and Free Riding*. 2004. doi.org ↗
- 26 VulnCheck. *State of Exploitation 2026*. 2026. vulncheck.com ↗
- 27 HIPAA Journal. *Patient Death Linked to Ransomware Attack on NHS Pathology Provider*. 2025-06. hipaajournal.com ↗
- 28 InfoWorld. *Internet Bug Bounty program hits pause on payouts*. 2026-04-03. info-world.com ↗
- 29 Sky TG24; Il Fatto Quotidiano. *Truffa con nome (e voce) del ministro Crosetto: cosa sappiamo*. 2025-02-08. tg24.sky.it ↗
- 30 UK Finance. *Annual Fraud Report 2026 (press release)*. 2026-06. ukfinance.org.uk ↗
- 31 K. Hackenburg, B. M. Tappin, L. Hewitt, E. Saunders, S. Black, H. Lin, C. Fist, H. Margetts, D. G. Rand y C. Summerfield. *The levers of political persuasion with conversational artificial intelligence*. 2025-12. science.org ↗
- 32 M. Spence. *Job Market Signaling*. 1973. doi.org ↗
- 33 C. Schmitz, L. Hammond y A. Chan. *Characterizing Agentic Flooding of Government Services*. 2026-08-17. arxiv.org ↗
- 34 The Decoder (con el Financial Times). *AI fraud forces world's largest accounting body to end online exams*. 2025-12-29. the-decoder.com ↗
- 35 Naciones Unidas. *General Assembly Adopts More Than 60 Resolutions, Decisions of Its First Committee (GA/12736)*. 2025-12-01. press.un.org ↗
- 36 Tom's Hardware. *China's chip champions ramp up production of AI accelerators at*

- domestic fabs, but HBM and fab production capacity are towering bottlenecks*. 2026. tomshardware.com ↗
- 37 D. Hendrycks, E. Schmidt y A. Wang. *Superintelligence Strategy: Expert Version*. 2025-03-07. arxiv.org ↗
 - 38 Oklahoma Watch. *Have 70 % of casualties in the Russia-Ukraine war been caused by drones?*. 2026-06-09. oklahomawatch.org ↗
 - 39 Bueno de Mesquita, Smith, Siverson y Morrow. *The Logic of Political Survival*. 2003. doi.org ↗
 - 40 Drago y Laine. *The Intelligence Curse*. 2025-04. intelligence-curse.ai ↗
 - 41 J. R. Arnold (AI Frontiers). *Superintelligence Deterrence Has an Observability Problem*. 2025-08-14. ai-frontiers.org ↗
 - 42 METR. *Investigation of the OpenAI Hugging Face incident*. 2026-08-26. metr.org ↗
 - 43 J. Carlsmith. *Superforecasting the premises in «Is power-seeking AI an existential risk?»*. 2023-10-18. joecarlsmith.com ↗
 - 44 G. Kestin y N. Soares. *Views on AI Existential Risk Before and After a Public Event at Harvard University*. 2026-03-29. arxiv.org ↗
 - 45 M. Strathern. *‘Improving ratings’: audit in the British University system*. 1997. cambridge.org ↗
 - 46 Anthropic. *Agentic Misalignment in Summer 2026*. 2026. alignment.anthropic.com ↗
 - 47 LawZero. *Yoshua Bengio launches LawZero: a new nonprofit advancing safe-by-design AI*. 2025-06-03. lawzero.org ↗
 - 48 J. Carlsmith. *Is Power-Seeking AI an Existential Risk?*. 2022-06-16. arxiv.org ↗
 - 49 TechCrunch (R. Bellan). *«Gambling with our lives»: Anthropic researcher quits, warns against self-improving AI*. 2026-09-09. techcrunch.com ↗
 - 50 Bloomberg (M. Barkley). *AI Pioneer Andrew Ng Calls Extinction Fears «Science Fiction»*. 2026-09-17. bloomberg.com ↗
 - 51 Kulveit et al.. *Gradual Disempowerment: Systemic Existential Risks from Incremental AI Development*. 2025-01. arxiv.org ↗
 - 52 Meyer et al. (Microsoft). *How effective is multifactor authentication at deterring cyberattacks?*. 2023-05. arxiv.org ↗
 - 53 Cornish y Clarke. *Opportunities, Precipitators and Criminal Decisions: A Reply to Wortley’s Critique of Situational Crime Prevention*. 2003. popcenter.asu.edu ↗
 - 54 OpenAI. *Strengthening ChatGPT’s responses in sensitive conversations*. 2025-10-27. openai.com ↗
 - 55 UNICEF y Grupo Interinstitucional de la ONU para la Estimación de la Mortalidad en la Niñez. *Progress in reducing child deaths slows as 4.9 million children under five die in 2024*. 2026-03-18. unicef.org ↗