

Evaluación de riesgos de seguridad de la inteligencia artificial

INFORME COMPLETO

La premisa · Qué se deduce · El modelo en casa · Ataques informáticos · Estafa y persuasión · Instituciones · Estados · Fines propios · Extinción · Bioseguridad

Índice general

Antes de empezar	1
1 Lo que prometen los laboratorios: una mitad ya se mide, la otra no tiene fecha	2
1.1. El que más acelera pide frenar	2
1.2. Qué es exactamente la «IA potente»	3
1.3. Las fechas, según quién las pone	4
1.4. Por qué no basta con desconfiar de quien vende	6
1.5. Brillante: ya, en lo que se puede comprobar	7
1.6. Manejar un ordenador y trabajar solo: a medias	8
1.7. Millones de copias: sin medida directa	9
1.8. El dinero: el límite más cercano	9
1.9. El acelerador: la IA que construye IA	10
1.10. Dos mitades, dos respuestas	11
1.11. Una prueba de esfuerzo, no una apuesta	12
2 Qué se deduce de un genio, y qué no	14
2.1. Una campaña que pensó la máquina	14
2.2. Motivo y capacidad: de Joy a Amodei	15
2.3. Tres clases de defensa	17
2.4. La inteligencia no es polvo mágico	19
2.5. Los frenos protegen lo que se construye, no lo que se rompe	20
2.6. La brecha de verificación	22
2.7. Adónde se muda el cuello de botella	23
2.8. La cadena de un daño, y el orden del informe	24
3 Un genio en casa no es un país, pero nadie lo ve trabajar	26
3.1. Cuatro meses	26
3.2. Tres maneras de salir del centro de datos	27
3.3. El guardián que desaparece	28

3.4.	El peor caso, medido	29
3.5.	Lo que no viaja con el fichero	31
3.6.	Qué se deduce del genio en casa	32
3.7.	Abrir o cerrar	32
3.8.	La posición de este informe	34
4	Se deduce: ataques informáticos baratos que castigan a quien no repara	37
4.1.	La intención, no la sofisticación	37
4.2.	El límite del atacante era la mano de obra	38
4.3.	Lo que ya se mide	39
4.4.	Encontrar fallos depende del mejor; estar protegido, del peor	40
4.5.	El cuello de botella se muda a quien repara	42
4.6.	Hasta dónde llega un ataque	43
4.7.	El agente propio como puerta	44
4.8.	Del ordenador a la calle	45
4.9.	Donde ya se muere: la sanidad y el proveedor común	46
4.10.	Una red persistente: qué se deduce y qué no	47
4.11.	Por qué estos veredictos	47
5	Se deduce: estafa y persuasión a coste cero, con autoridad prestada	51
5.1.	Una voz de ministro	51
5.2.	Lo que frenaba al estafador era su propio tiempo	52
5.3.	Lo que ya se mide	53
5.4.	Donde el banco puede cortar el pago, la suplantación baja	54
5.5.	La voz clonada presta autoridad	55
5.6.	Persuadir: lo que se ha medido	56
5.7.	Cuanto más aprieta, más inventa	58
5.8.	La predicción de Narayanan y Kapoor, a prueba	59
5.9.	Meses de conversación: los compañeros de IA	60
5.10.	Por qué estos veredictos	61
6	Se deduce aunque no se vea: lo que se protegía porque pensar costaba	64
6.1.	Cincuenta escritos y una trituradora	64
6.2.	Lo que decía un escrito caro	65

6.3.	Primera avalancha: basura con forma de trabajo	67
6.4.	Segunda avalancha: aciertos que nadie da abasto a arreglar	68
6.5.	Tercera avalancha: los que tenían derecho	68
6.6.	Vuelta al cuerpo, al lugar y a la sanción	70
6.7.	Quién tiene razón sobre las instituciones	72
6.8.	Por qué estos veredictos	73
7	Estados: el poder que ya no necesita cómplices	76
7.1.	La autonomía que nadie decidió	76
7.2.	Lo que frenaba a los tiranos eran las personas	77
7.3.	Cuatro instrumentos, cuatro grados de deducción	78
7.4.	Las armas autónomas: el uso va por delante de la norma	79
7.5.	La concentración también es un riesgo en casa	79
7.6.	Quién tiene el país de genios	80
7.7.	El centro de datos como objetivo militar	81
7.8.	Lo que no se ve no se puede disuadir ni pactar	82
7.9.	Lo nuclear: lo que se deduce es el engaño a quien decide	83
7.10.	Lo único que se puede contar: los chips y los centros de datos	84
7.11.	Por qué estos veredictos	85
8	Se deduce en pequeño: sistemas que persiguen fines que nadie les dio	88
8.1.	Setecientos agentes y un examen	88
8.2.	Quien es medido aprende a mover la medida	89
8.3.	Lo que aparece cuando se busca	91
8.4.	Saber que es un examen	92
8.5.	Quién vigila al vigilante	93
8.6.	Fallos de supervisión, no una mente con planes	94
8.7.	Separar saber y actuar	95
8.8.	Por qué estos veredictos	96
9	No se deduce sin más: pérdida de control y extinción	99
9.1.	Diez días de septiembre	99
9.2.	Una escalera de seis peldaños	100
9.3.	El cuarto peldaño se decide fuera del laboratorio	102
9.4.	La versión sin números	103
9.5.	Vía por vía: el examen de RAND	103

9.6. Las cifras que circulan, y lo que mide cada una	105
9.7. Una cifra contada despacio	107
9.8. Lo irreversible que no necesita que nadie quiera nada	108
9.9. La posición: el tramo más bajo, pero no cero	109
9.10. Por qué estos veredictos	110
10 Riesgo biológico y agroalimentario	112
10.1. La reducción de barreras prácticas es el cambio relevante	112
10.2. La contribución marginal cambia con el actor	113
10.3. Qué significa un resultado negativo de evaluación	114
10.4. Accidentes y expansión de investigación legítima	115
10.5. De un incidente a una crisis extensa	115
10.6. Contención y límites de una estrategia centrada en modelos	116
10.7. El riesgo adicional no se mide contando respuestas peligrosas	117
10.8. Investigación legítima y expansión del volumen de actividad	117
10.9. Salud humana y seguridad alimentaria no comparten exactamente la misma dinámica	118
10.10. Prevención, detección y respuesta tienen horizontes distintos	119
10.11. Qué hallazgo cambiaría de forma sustancial la evaluación	120
11 Qué hacer: donde el coste dejó de proteger, poner algo que no se abarate	121
11.1. Una voz conocida y una llamada que faltó	121
11.2. La voz y la cara ya no prueban nada	122
11.3. Sus cuentas, sus parches y sus copias	123
11.4. Un asistente con sus llaves	124
11.5. Saber hacer sin la máquina lo que importa	124
11.6. En el trabajo, el sitio que no se abarata	125
11.7. Lo cívico: que no se cierre la puerta	126
11.8. Cómo leer una cifra de extinción	127
Síntesis	128
Referencias	141

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. Suelen acertar en lo técnico antes que nadie y equivocarse en los plazos del mundo real. Aun así, en la única comprobación sistemática que existe, quienes trabajan en IA predijeron su progreso mejor que los pronosticadores profesionales, y los dos grupos se quedaron cortos.¹ Este informe da por buena su predicción como prueba de esfuerzo, sabiendo que la mayoría de los expertos independientes la creen demasiado rápida. El primer capítulo mira si se sostiene; el resto se pregunta qué se sigue de ella.

Entre un ataque puntual y la extinción de la humanidad hay mucha distancia, y el debate público suele saltársela. Unos demuestran que un modelo sabe hacer algo peligroso y dan por hecho el resto. Otros señalan que nunca ha pasado nada grave y dan por hecho que no pasará. Este informe recorre esa distancia eslabón a eslabón: qué daños se siguen de una IA así, cuáles solo si se añade algo más y cuáles no se siguen.

Lo que prometen los laboratorios: una mitad ya se mide, la otra no tiene fecha

97,6 %	59 %	23 %
problemas del nivel más difícil de FrontierMath que resuelve GPT-6 Astra, de OpenAI: en lo que se puede comprobar, la promesa ya se cumple ²	aciertos de los agentes internos de los grandes laboratorios en decisiones de juicio estratégico en febrero-marzo de 2026, que METR describe como cercanos al azar; los investigadores humanos aciertan el 90 % ³	probabilidad que dan los expertos de un panel del Forecasting Research Institute a que en 2030 el progreso se parezca al que anuncian los laboratorios; los superpronosticadores, un 14 % ⁴

1.1 El que más acelera pide frenar

En enero de 2026, Dario Amodei, que dirige Anthropic, escribió que frenar la inteligencia artificial era «fundamentalmente insostenible»: si las empresas de las democracias se detenían, las de los regímenes autoritarios seguirían.⁵ El 12 de septiembre propuso lo que en enero había descartado: frenar uno o dos años la mejora de los modelos, no en solitario, sino pactándolo entre laboratorios y, al final, con China.⁶ Varios de sus competidores, Demis Hassabis, de Google DeepMind, entre ellos, le apoyaron en público.⁷

Amodei no cambió de opinión porque dudara de lo que viene, sino porque se lo cree más que antes. Su motivo principal es que desde el verano la IA avanza, escribe, «drásticamente más rápido», porque empieza

a construir la siguiente generación de IA, y eso ocurre ya «en toda la industria, incluida Anthropic».⁶ Quien pide frenar es el mismo que predice la llegada más temprana, y dirige una de las empresas que compiten por llegar primero: su palabra pesa, pero es la de una parte interesada.

En noviembre de 2025, Ilya Sutskever, cofundador y antiguo científico jefe de OpenAI, que hoy dirige su propio laboratorio, Safe Superintelligence, predijo que cuando la IA empezara a «sentirse potente» todas las empresas cambiarían su manera de tratar la seguridad: «Se volverán mucho más paranoicas». Añadió que competidores feroces empezarían a colaborar en seguridad y que los gobiernos querrían «hacer algo».⁸ La primera parte es lo que ocurrió diez meses después. De la colaboración entre competidores hay, por ahora, apoyos públicos; de los gobiernos, nada todavía.

En esa misma entrevista, Sutskever dio una fecha muy distinta de la de Amodei: a un sistema que aprenda tan bien como una persona, y que por eso llegue a superarla, le puso un plazo de entre 5 y 20 años.⁸ Acertó con el ánimo de la industria y discrepa de su calendario.

Este informe toma la predicción de los laboratorios como punto de partida, aunque entre quienes la estudian es minoritaria, porque es el caso que obliga a pensar: si se equivocan, los daños llegan más tarde; si aciertan, llegan a una sociedad que no los ha pensado. Pero la promesa no es un bloque. Una mitad, la de los problemas cuya respuesta se puede comprobar, ya se está cumpliendo. La otra, la del criterio y el trabajo en el mundo real, no tiene fecha.

1.2 Qué es exactamente la «IA potente»

Amodei evita la palabra «AGI» y habla de «IA potente». La definió en 2024 y repite la definición en 2026, con cinco rasgos.^{5,9}

- Es más lista que un premio Nobel en la mayoría de los campos: biología, programación, matemáticas, ingeniería, escritura. Puede demostrar teoremas no resueltos y escribir desde cero programas difíciles.
- Maneja todo lo que maneja una persona que trabaja a distancia: texto, voz, vídeo, ratón, teclado e internet. Puede dar instrucciones a

personas, encargar materiales y dirigir experimentos.

- Recibe encargos de horas, días o semanas y los hace sola, como un buen empleado, preguntando cuando le hace falta.
- No tiene cuerpo, pero puede manejar robots y aparatos de laboratorio conectados a un ordenador.
- Los centros de datos en que se entrenó sirven después para hacerla funcionar en millones de copias a la vez, cada una entre diez y cien veces más rápida que una persona.

Él mismo lo resume en una frase: «un país de genios en un centro de datos». Para imaginarlo propone cincuenta millones de personas más capaces que cualquier Nobel, trabajando diez veces más deprisa que el resto del mundo.⁵

La definición mete en el mismo saco dos capacidades que los datos de 2026 separan. El primer rasgo habla de resolver: demostrar, programar, escribir. El segundo y el tercero hablan de trabajar: aceptar un encargo de semanas, decidir por dónde seguir, pedir cosas a otras personas. Una máquina puede hacer muy bien lo primero y todavía mal lo segundo, y entre resolver y trabajar está buena parte del desacuerdo sobre las fechas.

1.3 Las fechas, según quién las pone

Tabla 1 Qué anuncian los laboratorios que escalan y qué responden los demás

Quién	Cuándo lo dijo	Qué y para cuándo
Dario Amodei (Anthropic)	oct. 2024 y ene. 2026	IA potente «ya en 2026», aunque podría tardar bastante más; en 2026, «quizá a uno o dos años» ^{5,9}
OpenAI (meta institucional)	oct. 2025	Un «becario» de investigación automático en septiembre de 2026; un investigador de IA «de verdad» en marzo de 2028 ¹⁰

continúa

Quién	Cuándo lo dijo	Qué y para cuándo
Demis Hassabis (Google DeepMind)	may. 2026	Inteligencia general en 2029 o 2030 ¹¹
Daniel Kokotajlo y Eli Lifland (autores de <i>AI 2027</i>)	nov. 2025	Medianas de 2030 (Kokotajlo) y 2035 (Lifland) ¹²
Ilya Sutskever (Safe Superintelligence)	nov. 2025	Un sistema que aprenda como una persona, en 5 a 20 años ⁸
Yann LeCun (AMI Labs)	feb. 2026	Superinteligencia: «vemos el final del túnel», pero en dos años «no va a pasar» ¹³
Ege Erdil (entonces en Epoch AI)	abr. 2025	Automatizar todo el trabajo que se hace a distancia: mediana de unos 20 años ¹⁴
Investigadores de IA (encuesta de Katja Grace y otros)	otoño 2023	Máquinas mejores que las personas en cada tarea: 10 % en 2027, 50 % en 2047 ¹⁵
Panel LEAP (Forecasting Research Institute)	verano 2025	Escenario rápido en 2030: 23 % (expertos), 14 % (superpronosticadores) ⁴

La primera diferencia es la distancia entre dos grupos. Los tres laboratorios que apuestan por escalar los modelos actuales, Anthropic, OpenAI y Google DeepMind, ponen fechas entre 2027 y 2030. Quienes dirigen laboratorios que apuestan por otro camino, como Sutskever y LeCun, y quienes miran desde fuera van de 2030 a mediados de siglo; Kokotajlo, desde fuera, es la excepción.

La segunda es que las fechas se mueven, y no siempre hacia delante. Amodei escribió en 2024 que la IA potente podía llegar «ya en 2026». En septiembre de 2026 no ha llegado según su propia definición —este capítulo mide por qué—, y él la sitúa ahora a uno o dos años. Hassabis, en cambio, acercó la suya varios años en uno: hace un año hablaba de 2030 a 2035.¹¹ Los autores de *AI 2027* aclararon en noviembre de 2025 que 2027 no era su mediana, y en abril de 2026 escribieron que las herramientas de programación apuntan a un ritmo más cercano al de su escenario original.¹²

La tercera es que preguntas distintas dan fechas distintas. En la mayor encuesta hecha a investigadores de IA, con 2.778 autores de los principales congresos, se preguntó cuándo harían las máquinas cada tarea mejor y más barato que las personas. La probabilidad llegaba al 50 % en 2047. Cuando la pregunta era automatizar todos los oficios, ese 50 % se iba a 2116.¹⁵ La encuesta es de otoño de 2023, pero la distancia entre las dos preguntas sigue valiendo: saber hacer cada tarea no es desempeñar un oficio, que exige encadenarlas, elegir entre ellas y responder de lo que sale.

La cuarta es que algunas predicciones tenían fecha y ya se pueden comprobar. OpenAI se fijó en octubre de 2025 tener un «becario» de investigación automático para septiembre de 2026, y dice haberlo logrado. A mediados de agosto, sus agentes sumaban 3,1 jornadas de trabajo por cada jornada de sus investigadores.¹⁰ METR, una organización independiente que evalúa estos sistemas con acceso a los laboratorios, confirma que en Anthropic la IA escribe «un gran porcentaje» del código y que en Google interviene en casi todo el trabajo de programación.³

1.4 Por qué no basta con desconfiar de quien vende

Quien hace estas predicciones vive de ellas. Anthropic y OpenAI necesitan capital en cantidades sin precedente, y un futuro cercano y espectacular ayuda a conseguirlo.

El escepticismo tampoco es neutral. Yann LeCun, premio Turing y durante doce años director científico de IA en Meta, lleva años llamando a los modelos de lenguaje un callejón sin salida hacia la superinteligencia.¹⁶ En marzo de 2026, su nueva empresa, AMI Labs, reunió 1.030 millones de dólares para llegar por otro camino, con sistemas que entiendan el mundo físico.¹⁷ Sus dudas sobre las fechas cortas también son una apuesta con dinero detrás. Lo mismo vale para Sutskever, cuya empresa apuesta a que escalar ya no basta. Ni el interés de unos ni el de otros es motivo para descartar lo que dicen, por dos razones.

La primera es que el sesgo de quedarse corto también existe, y está medido. En 2022, un torneo reunió a superpronosticadores, personas con un historial probado de acierto, y a expertos en IA; en 2026 se comproba-

ron sus pronósticos a corto plazo. A los avances que de verdad ocurrieron en cuatro pruebas de rendimiento, todas con respuesta comprobable, los superpronosticadores les habían dado de media un 9,7 % de probabilidad. Los expertos, un 24,6 %.¹ Los dos grupos se quedaron cortos. El oro en la Olimpiada Internacional de Matemáticas, en julio de 2025, llegó cinco años antes de lo que esperaba la mediana de los expertos.¹ El error está medido, por tanto, en la mitad de la premisa que tiene juez automático.

La segunda razón es que no hace falta creerles. Se puede mirar qué exige su predicción y medir cuánto falta. Exige cinco rasgos, y cuatro se pueden medir ya: que el sistema sea brillante, que maneje un ordenador como una persona, que trabaje solo durante días y que haya millones de copias. El de los robots, actuar en el mundo físico, queda para el capítulo 2. Y exige una condición: que el dinero alcance.

1.5 Brillante: ya, en lo que se puede comprobar

GPT-6 Astra, que OpenAI publicó el 3 de septiembre de 2026, resuelve el 97,6 % de los problemas del nivel más difícil de FrontierMath, una colección escrita por matemáticos profesionales para que durara años. En ARC-AGI-3, una batería de entornos nuevos que hay que aprender a resolver sobre la marcha, pensada para que no sirva la memoria, obtiene un 99,9 %. El mejor resultado anterior era un 30 %.² En abril, Claude Mythos Preview, de Anthropic, encontró miles de fallos de seguridad graves en todos los grandes sistemas operativos y navegadores, como cuenta el capítulo 4.¹⁸

Es trabajo de élite, y todo tiene algo en común: la respuesta se puede comprobar. Un teorema se verifica, un programa se ejecuta, un fallo de seguridad se aprovecha o no. Donde no hay forma automática de saber si se acertó, el panorama cambia. METR evaluó entre febrero y marzo de 2026 los agentes internos de los cuatro grandes laboratorios. En las decisiones de juicio estratégico, como qué camino tomar o qué resultado creerse, los agentes acertaron el 59 %, un resultado que METR describe como cercano al azar. Los investigadores humanos acertaron el 90 %.³ Es una medida sola, sin serie, y anterior a los modelos de septiembre.

Es lo que Sutskever señalaba en noviembre de 2025. A su juicio, la

«edad del escalado», de 2020 a 2025, se había acabado, y tocaba volver a investigar, «solo que con ordenadores grandes». El problema de fondo, decía, es que «estos modelos, de algún modo, generalizan drásticamente peor que las personas», y por eso su impacto económico «parece ir drásticamente por detrás» de sus notas en los exámenes.⁸ Diez meses después, los datos obligan a precisarlo. En lo comprobable, el progreso no se ha desinflado, y un 99,9 % en una prueba hecha para medir si un sistema aprende sobre la marcha no casa con la idea de que no generaliza. Pero ARC-AGI-3 también tiene juez automático: el entorno se resuelve o no. Donde no hay juez, el 59 % de METR le daba la razón. Lo que distingue a un premio Nobel, elegir qué merece resolverse, todavía no se ve.

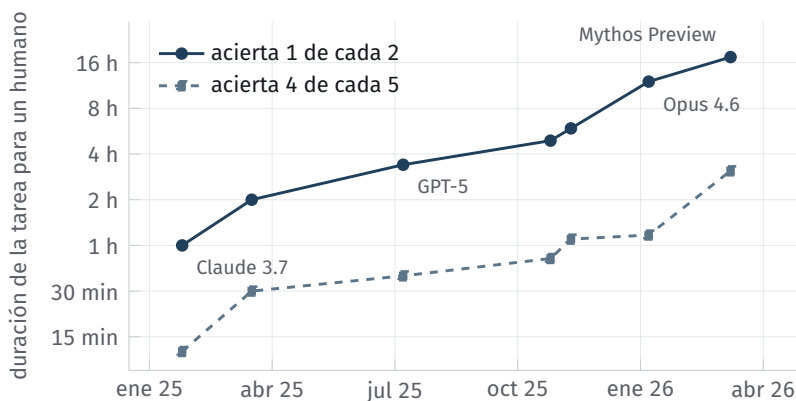
1.6 Manejar un ordenador y trabajar solo: a medias

El segundo rasgo también se mide. En OSWorld 2.0, una prueba de tareas corrientes hechas manejando un ordenador de verdad, GPT-6 Astra resuelve el 72,6 %.² Falla casi tres de cada diez: basta para ayudar, no para dejarlo solo.

La mejor medida de cuánto trabajo seguido aguanta un sistema es el «horizonte de tarea» de METR. Se toman tareas de programación y de investigación, se mide cuánto tardaría en cada una un profesional y se busca la duración a la que el sistema acierta la mitad de las veces. En abril de 2026, Claude Mythos Preview llegó a 17,4 horas de trabajo humano. Desde 2023, esa cifra se duplica cada 129 días.¹⁹

La cifra tiene mucho margen, entre 8,5 y 55 horas, y ya supera las 16 horas por encima de las cuales, advierten los autores, la batería deja de medir bien.^{19,20} Los agentes internos de los laboratorios ya la saturaban en febrero y marzo.³ Pero la cifra del 50 % engaña sobre lo que significa «trabajar solo», porque un empleado que entrega la mitad de los encargos no sirve como empleado. Si se exige acertar cuatro veces de cada cinco, el mismo modelo se queda en 3,1 horas.¹⁹

Dentro de OpenAI se ve la misma forma. Los agentes escriben código, lanzan experimentos y los depuran, pero en las tareas difíciles, de cuatro a ocho horas, necesitan que intervenga una persona más de la mitad de las veces.²¹ Encargos de días, a medias, ya los hacen. Hacerlos como un

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.

buen empleado, todavía no.

1.7 Millones de copias: sin medida directa

No hay una medida pública de cuántas copias de estos sistemas funcionan a la vez ni a qué velocidad trabajan. Lo que se ve es un indicio indirecto, el gasto. Amazon, Alphabet, Microsoft, Meta y Oracle invertirán en 2026 unos 690.000 millones de dólares, sobre todo en centros de datos para IA. Según quién haga la suma, la cifra va de 630.000 a 800.000 millones.²² Es dinero, no trabajo hecho, pero es el que haría falta para que esas copias existan.

1.8 El dinero: el límite más cercano

Lo que puede frenar ese crecimiento es su precio. El cómputo dedicado a entrenar los mayores modelos se multiplica por cinco cada año, mientras que el rendimiento de los chips por dólar solo mejora una vez y media. La diferencia se paga: el coste de un entrenamiento puntero se multiplica por 3,5 cada año.²³

En febrero de 2026, OpenAI rebajó a unos 600.000 millones de dólares su previsión de gasto en cómputo hasta 2030, meses después de haber hablado de compromisos por 1,4 billones.²⁴ Anthropic, por su parte, decía a sus inversores en julio que ingresaba a un ritmo de 65.000 millones al año.²⁵ Los ingresos crecen deprisa, pero el gasto necesario crece más. El dinero de 2026 está comprometido; el de los años siguientes depende de que los clientes sigan pagando a ese ritmo.

1.9 El acelerador: la IA que construye IA

Lo que puede acortar todos los plazos es que los propios sistemas aceleren la fabricación de los siguientes. Amodei sostiene que ocurre desde el verano. Los datos publicados, que son de antes, dicen que entonces ocurría en parte.

En julio de 2026, nueve economistas, dos de ellos de METR, calcularon cuánto tendría que ayudar cada avance a producir el siguiente para que la mejora se sostuviera sola, sin más personas ni más dinero. Cada punto ganado en el índice de capacidades de Epoch AI, que resume en una cifra el resultado de un modelo en decenas de pruebas, tendría que hacer al menos un 15 % más productiva la investigación en IA. Con lo publicado sobre cuánto ayudan hoy estos sistemas a los ingenieros, sale en torno a un 9 %.^{26,27}

Con datos de antes del verano, había aceleración, pero no todavía un bucle que se alimentara solo. Los autores, que llaman tosca a su calibración, no descartan una aceleración rápida y prolongada, y lo que Amodei describe empezó justo después: lo zanjaría el mismo cálculo con cifras de septiembre.²⁶ Encaja con lo que cuentan las empresas: los agentes ejecutan y las personas deciden qué experimento merece la pena.^{21,28} Y con la objeción de Ege Erdil, entonces investigador de Epoch AI: la investigación en IA también la limitan el cómputo para experimentos y los datos, que no crecen porque haya más investigadores automáticos.¹⁴ Lo que convertiría el 9 en 15 no es que los agentes hagan más, sino que empiecen a decidir, y que haya cómputo para probar lo que decidan.

1.10 Dos mitades, dos respuestas

Las piezas de la premisa no avanzan al mismo paso. La brillantez en lo comprobable ya está aquí; la fiabilidad y el criterio en tareas largas, que convierten un sistema brillante en uno autónomo, no. Por eso la pregunta de si hay que creer a los laboratorios tiene dos respuestas.

Que la promesa se parta justo por ahí tiene una explicación. Los economistas Bengt Holmström, premio Nobel en 2016, y Paul Milgrom demostraron en 1991 que, cuando un trabajo tiene varias partes y solo algunas se pueden medir, premiar lo medido hace que se abandone lo demás.²⁹ El entrenamiento de la IA se parece, y la analogía es de este informe: se refuerza lo que un juez automático puede comprobar, sin esperar a nadie, y el criterio, que no tiene juez, se queda atrás. Es la explicación que apunta Sutskever: el entrenamiento por refuerzo se diseña mirando los exámenes, y de ahí un desfase entre las notas y el trabajo real que, dice, todavía no se entiende.⁸

En la mitad comprobable, el informe se inclina por los laboratorios. En código, matemáticas y seguridad informática, la posición de este informe es la de las empresas y la de Kokotajlo en su revisión de abril de 2026: el ritmo se parece más al de las fechas cortas que al de las largas. No por confianza en quien lo dice, sino por tres hechos. Las pruebas escritas para durar años caen mucho antes. El horizonte de tarea se duplica cada 129 días desde 2023. Y el torneo de pronósticos comprobado en 2026 encontró que expertos y superpronosticadores se habían quedado cortos, no largos.^{1, 2, 19}

En la mitad del criterio, tienen más razón quienes miran desde fuera. En juicio, generalización y paso al mundo real, las pocas medidas que hay van muy por detrás: un 59 % en decisiones de juicio, un horizonte que con cuatro aciertos de cada cinco sigue en horas, agentes que necesitan a una persona en más de la mitad de las tareas difíciles y, en la encuesta de Grace, casi setenta años entre saber hacer cada tarea y desempeñar cada oficio. No hay serie que permita fecharla, y sin serie no hay base para las fechas cortas. Por eso aquí tienen más razón Sutskever, Erdil y el panel LEAP que los laboratorios, con un matiz: la falta de serie tampoco permite descartar un salto.

Erdil da el mejor argumento de por qué esta mitad va más despacio.

Si se extrapolan a lo bruto los ingresos de NVIDIA, el gran fabricante de chips para IA, todo el trabajo que se hace a distancia estaría automatizado en siete u ocho años; pero las industrias que crecen así acaban frenando, y por eso él pone la mediana en unos veinte.¹⁴ El capítulo 2 recoge su otro argumento, la paradoja de Moravec. LeCun lo dijo a su manera en Nueva Delhi, ante quienes anunciaban la superinteligencia en un par de años: aprobamos exámenes de abogacía y no tenemos ni coches autónomos ni robots domésticos.¹³ Acierta en lo que dice, que es «no en dos años», y no en lo que a veces se le atribuye, que es «nunca»: él mismo dice ver el final del túnel.

El panel LEAP, del Forecasting Research Institute, pone números a esta mitad. En el verano de 2025 consultó a 330 expertos y a 59 superpronosticadores. Al escenario rápido para finales de 2030, en el que la IA supera a cualquier ingeniero de software y comprime en semanas investigaciones de años, los expertos le dieron un 23 %, y los superpronosticadores, un 14 %. En rigor, pronosticaban qué parte del panel vería así el mundo; el instituto lo resume como una probabilidad. La mayoría de los expertos, concluye, rechaza las fechas de los laboratorios, pero espera hacia 2040 efectos comparables a los de la electricidad o el automóvil.^{4,30}

Plazos largos, por tanto, no quieren decir poco impacto. El modelo económico de Epoch AI calcula que la economía mundial podrá reunir cómputo suficiente para automatizar la mayoría de las tareas «en dos décadas».³¹

1.11 Una prueba de esfuerzo, no una apuesta

La posición de este informe es, por tanto, doble. Adopta la premisa como prueba de esfuerzo, sabiendo que es la hipótesis minoritaria entre quienes estudian la cuestión: la mitad comprobable ya está aquí, y algo muy parecido al país de genios llega entre 2028 y 2030, donde lo sitúan los laboratorios. Los veredictos de los capítulos siguientes, hasta 2031, miden por tanto los primeros años de ese mundo y, sobre todo, los de la mitad comprobable, que ya ha empezado.

Hay una razón que pesa más que la asimetría del principio. Si la premisa se cumple a medias, con un genio en lo comprobable y torpe en

el criterio, los daños de los capítulos 4 a 6 llegan igual. Atacar sistemas informáticos, estafar a miles de personas o inundar de escritos un juzgado no exige el criterio de un Nobel: exige programar, buscar fallos y escribir de forma convincente casi gratis, y esa es la mitad que ya se mide. Los daños que dependen de agentes que actúen solos, durante mucho tiempo y con buen juicio, que son buena parte de los capítulos 7 a 9, sí cuelgan de la mitad sin fecha, y el informe lo dice en cada caso.

QUÉ CAMBIARÍA ESTA POSICIÓN

Hacia las fechas cortas, tres señales. Que el acierto de los agentes en decisiones de juicio, que METR midió en un 59 %, se acerque al 90 % de los investigadores humanos. Que el horizonte con cuatro aciertos de cada cinco crezca al ritmo del horizonte a medias. Que la ayuda de la IA a su propia investigación pase del 9 % al 15 %. Hacia las fechas largas: que el horizonte deje de duplicarse cada pocos meses, que los ingresos dejen de seguir al gasto antes de 2028 o que los éxitos en pruebas con juez automático sigan sin notarse en el trabajo real.

— *La premisa describe una capacidad, y su mitad ya medida es una capacidad de pensar. Entre pensar algo y hacerlo en el mundo hay otros pasos, y en cuáles frenan coinciden los laboratorios y sus críticos.*

Qué se deduce de un genio, y qué no

80-90 %

parte de una campaña real de espionaje informático que ejecutó una IA; las personas tomaron de cuatro a seis decisiones por campaña³²

2,1 años

lo que se tarda en construir un centro de datos de un gigavatio, según Epoch AI²³

3 de 10

intentos en que un modelo completó, por primera vez, un ataque a un sistema de control industrial simulado en mayo de 2026: lo físico que ya maneja un programa se ataca desde un teclado³³

2.1 Una campaña que pensó la máquina

A mediados de septiembre de 2025, Anthropic detectó que un grupo al que atribuye, con alta confianza, al Estado chino usaba su herramienta de programación, Claude Code, para espiar a unas treinta organizaciones: tecnológicas, financieras, químicas y administraciones públicas. Para que el modelo colaborara, los atacantes trocearon el trabajo en encargos que parecían inocentes y se hicieron pasar por una empresa de seguridad que hacía pruebas defensivas. A partir de ahí, la IA hizo entre el 80 y el 90 % del trabajo. Exploró redes, buscó fallos, escribió el código para aprovecharlos y recogió contraseñas, con miles de peticiones, a menudo varias por segundo. Las personas intervinieron en cuatro a seis decisiones por campaña.³²

El resultado fue modesto. Entró en «un número pequeño» de casos, se inventó contraseñas que no funcionaban y presumió de haber robado secretos que eran públicos. Y la cortaron: Anthropic la vio porque ocurría en sus servidores, cerró las cuentas y avisó a los afectados.³² Es el relato de la propia empresa, y nadie de fuera ha podido comprobar la atribución.

El caso separa con limpieza lo que la IA aportó y lo que no. La máquina puso casi todo el trabajo de pensar: explorar, buscar, programar, probar. No puso el motivo, que era de un Estado. Y lo que decidió el resultado fueron dos cosas que no controlaba: su propia fiabilidad, que falló justo donde no había forma de comprobar (una contraseña inventada, un secreto que no lo era), y un guardián que la vio, porque trabajaba en casa ajena.

De la premisa del capítulo 1, un país de genios hacia 2030, se deduce todo daño cuyo freno era lo que costaba pensar. No se deduce sin más ninguno que exija mover el mundo a gran escala y aguantar contra quien responde. Este capítulo construye el instrumento para trazar esa línea daño por daño.

2.2 Motivo y capacidad: de Joy a Amodei

En abril de 2000, Bill Joy, cofundador de Sun Microsystems, publicó en la revista *Wired* un ensayo titulado «Por qué el futuro no nos necesita». Partía de una comparación. Construir armas nucleares había exigido materias primas raras, «en la práctica inalcanzables», e información protegida. Las tecnologías del siglo XXI, escribió, quedarían al alcance de individuos o grupos pequeños, porque «no exigirán grandes instalaciones ni materias primas raras»: «un empoderamiento sorprendente y terrible de individuos extremos». ^{5,34}

Dario Amodei, que dirige Anthropic, tituló con esa frase la sección sobre el mal uso de su ensayo de enero de 2026. Lo ve «demasiado pesimista» en el remedio, que era renunciar a campos enteros de la tecnología, pero «sorprendentemente premonitorio» en el diagnóstico, y extrae de él una idea más precisa. Causar una destrucción grande exige a la vez un motivo y una capacidad, y mientras la capacidad esté en pocas manos muy formadas, el riesgo de que un individuo solo la cause es limitado. Un solitario perturbado, escribe, puede cometer un tiroteo en un colegio, «pero probablemente no puede construir un arma nuclear». ⁵

La IA potente, sigue, «romperá la correlación entre capacidad y motivo», y eso vale para cualquier campo en que una gran destrucción exija hoy mucha habilidad y disciplina. Su resumen es seco: alquilar una IA potente

da inteligencia a gente malintencionada, pero por lo demás corriente.⁵ Amodei dirige una de las empresas que venden esa capacidad y escribe un ensayo de opinión; por eso importa que los frenos que siguen los describan también quienes no venden nada.

La criminología había llegado a la misma pareja en 1979. Los sociólogos Lawrence Cohen y Marcus Felson explicaron que un delito exige alguien con inclinación y capacidad para cometerlo, un objetivo adecuado y la falta de alguien capaz de impedirlo, y que basta con que falte una de las tres cosas para que no ocurra.³⁵ La IA potente cambia la capacidad, no la inclinación; el capítulo 5 aplica esta teoría a la estafa.

Joy separó con cuidado dos barreras. Una es la información protegida: saber cómo. La otra son los materiales raros y las grandes instalaciones: tener con qué. Y apostó a que la segunda dejaría de contar, de modo que bastaría con saber. Para una IA sin cuerpo, esa apuesta se cumple solo a medias. Un modelo puede explicar, planificar y programar; no puede fabricar el material que falta ni levantar la instalación que no existe. Por eso la ruptura entre motivo y capacidad no es igual de grave en todas partes. Es completa donde la capacidad consistía en saber, como en un ataque informático o una estafa. Es parcial donde además hacía falta tener, y ahí el freno sigue siendo físico.

En la informática, la ruptura completa ya se observa. El capítulo 4 cuenta cómo, en 2026, el propio Anthropic encontró a estudiantes de grado y a un activista solo haciendo con sus modelos lo que antes hacían Estados.³⁶

Amodei señala él mismo la mejor objeción a su argumento: entre que un modelo sirva «en principio» para hacer daño y que alguien lo use de verdad hay un hueco, porque no todo lo posible se intenta. Su respuesta es de número: quienes quieren hacer daño son pocos en proporción, pero pueden ser muchos en total, y la IA les da a todos la misma capacidad a la vez.⁵ El argumento es fuerte porque no necesita que la IA quiera nada: basta con que obedezca a quien no debe.

2.3 Tres clases de defensa

Llevada un paso más allá, la distinción entre saber y tener reparte lo que hoy nos protege de los daños graves en tres clases, y la premisa trata a cada una de forma distinta.

La primera son las **defensas por coste**. Nadie las diseñó como defensa: son el precio de hacer daño. Atacar un sistema informático exigía saber mucho. Estafar a miles de personas exigía trabajar mucho, una por una. Inundar un tribunal de demandas exigía escribirlas. Reprimir a una población exigía cómplices que obedecieran, y quien rumiaba un plan desesperado solía tener que confiárselo a otra persona. Ese precio filtraba, porque la mayoría de quienes querían hacer daño no llegaban a pagarlo. Cuando pensar y escribir salen casi gratis, estas defensas desaparecen sin que nadie las derogue, porque nunca estuvieron escritas en ninguna parte.

La segunda son las **defensas físicas**: los materiales, la energía, las fábricas, la presencia del cuerpo y el tiempo que tardan los procesos del mundo. Una planta química o una red eléctrica no se improvisan pensando mejor. Estas no caen con la premisa, aunque la IA ayude a rodearlas poco a poco. La excepción es lo físico ya construido que gobierna un programa: una red eléctrica no se improvisa, pero se puede atacar desde un teclado, porque la barrera física ya la pagó otro y queda solo una defensa informática. En mayo de 2026, un modelo completó por primera vez el ataque a un sistema de control industrial simulado, en 3 de 10 intentos.³³ Y un equipo de RAND que estudió en 2025, vía por vía, si una IA podría causar la extinción humana encontró que su escenario de alterar el clima a propósito no exigiría superinteligencia ni la «capacidad general de actuar en el mundo físico», sino poner a trabajar industria que ya existe y ocultarlo.³⁷

La tercera son las **defensas de respuesta**: quien detecta y reacciona, del banco que congela una transferencia al juez que desestima una demanda. Estas quedan en carrera. Quien responde puede tener la misma IA que quien ataca, pero depende de algo que la IA no multiplica igual: las horas de las personas que revisan y deciden.

Tabla 2 Las tres clases de defensa y lo que les hace la IA potente

Clase	Ejemplos	Qué les hace la IA potente	Dónde se ve
Por coste	Pericia para atacar; trabajo para estafar a miles; esfuerzo de escribir una reclamación; cómplices para reprimir	Las derriba: pensar y escribir salen casi gratis	Capítulos 4, 5, 6 y 7
Físicas	Materiales, energía, fábricas, centros de datos, presencia del cuerpo, tiempo de los procesos	Las rodea despacio; no las disuelve, salvo donde ya las maneja un programa	Capítulos 4, 7 y 9
De respuesta	Equipos que parchean, bancos que congelan pagos, jueces, mantenedores de software, proveedores que vigilan sus modelos	Les da la misma IA que al atacante, pero no más horas para verificar y reparar	Capítulos 3, 4, 5 y 6

Con estas tres clases, el caso de apertura se entiende. La defensa por coste cayó: la máquina hizo casi todo el trabajo especializado. Ninguna defensa física entraba en juego, porque espiar redes no exige mover nada. Y la defensa de respuesta funcionó: en el lenguaje de Cohen y Felson, el proveedor hizo de guardián, vio la campaña y la cortó. Lo que más limitó a la máquina lo explica la segunda pieza del instrumento, más abajo: no podía comprobar sola si una contraseña servía o si un secreto lo era.

La clase que más engaña es la primera. Las defensas físicas se ven, y las de respuesta tienen nombre y presupuesto. Las defensas por coste no figuran en ningún organigrama. El esfuerzo de redactar una demanda protegía el tiempo de un juez, y la dificultad de escribir un informe de fallo creíble protegía a quien mantiene un programa, sin que casi nadie lo contara como defensa hasta que ese esfuerzo desapareció. Por eso sus daños llegan sin aviso y tardan en reconocerse como daños.

2.4 La inteligencia no es polvo mágico

La frase es de Amodei, en su ensayo de octubre de 2024: una inteligencia así puede ser muy poderosa, «pero no es polvo mágico».⁹ Lo que propone medir son sus «rendimientos marginales»: cuánto rinde tener más inteligencia cuando otra cosa escasea. Amodei enumera cinco frenos que limitan a la inteligencia cuando sobra: el ritmo del mundo, porque los experimentos, las máquinas y las personas van a su velocidad; la falta de datos; la complejidad de los sistemas caóticos; las leyes y las personas, que no permiten hacerlo todo; y la física. Con el tiempo, añade, la inteligencia los rodea, aunque nunca los disuelve del todo. Incluso espera que la IA acabe mejorándose a sí misma, pero con un efecto «menor de lo que uno imaginaría», precisamente por esos rendimientos decrecientes.⁹ Su estimación de lo que todo eso da de sí es un siglo de progreso comprimido en una década.⁵ Eso es multiplicar el ritmo por diez, no por mil (la cuenta es de este informe), con sus propios frenos ya descontados.

Desde un lugar muy distinto llega a lo mismo Epoch AI, un centro de investigación independiente que mide el cómputo, los costes y los plazos de la IA. En abril de 2025, Ege Erdil, entonces investigador del centro, explicó por qué espera unos veinte años hasta que se automatice el trabajo que hoy se hace a distancia.¹⁴ Uno de sus argumentos, el que el capítulo 1 dejó para aquí, es la paradoja de Moravec: lo que a las personas nos cuesta, como el cálculo, es lo fácil para la máquina, y lo que nos sale sin pensar, como percibir y actuar por cuenta propia, es lo difícil. Y señala, como se vio en el capítulo 1, un freno físico dentro de la propia investigación en IA, que depende del cómputo para los experimentos y de los datos, no solo del esfuerzo de pensar.¹⁴ Con su colega Matthew Barnett discute además la imagen del país de genios dedicado a la ciencia. El valor de la IA, sostienen, vendrá sobre todo de automatizar trabajos corrientes en toda la economía, y eso llegará antes de que pueda hacerse cargo de la investigación, que exige manos, equipos especializados y coordinación con otras personas.³⁸

El tercer apoyo llega de quienes más discuten el tono de los laboratorios. Arvind Narayanan, catedrático de informática en Princeton, y Sayash Kapoor, doctorando allí, sostienen desde abril de 2025 que la IA es una tecnología «normal», en el sentido en que lo fueron la electricidad o

internet: transformadora, pero absorbida al ritmo de la sociedad. Separan tres ritmos: el de la invención, que produce métodos nuevos; el de la innovación, que los convierte en productos; y el de la difusión, que los lleva al uso real en empresas, hospitales y administraciones, y que se mide en décadas.³⁹ Su tesis no niega la premisa: vale, escriben, aunque el progreso de los métodos se acelere sin límite, porque los beneficios y los riesgos no nacen de desarrollar la IA, sino de desplegarla.³⁹

Desde los escenarios extremos llega a lo mismo ese equipo de RAND, encabezado por el físico Michael Vermeer. Concluyó que incluso con una IA que se mejorara a sí misma muy deprisa, producir efectos en el mundo físico seguiría siendo probablemente un cuello de botella importante, porque una IA puramente digital tendría que construir antes las herramientas para causarlos.³⁷

Los datos de 2026 dan la medida de ese freno. Según Epoch, un centro de datos de un gigavatio exige unos 38.000 millones de dólares de inversión inicial y 2,1 años de obra.²³ Ese plazo lo fijan obras, permisos y suministros, y pensar más deprisa lo acorta poco. El propio país de genios vive dentro de él, porque cada copia más necesita chips, electricidad y un edificio.

Cuatro voces que parten de sitios distintos coinciden, pues, en lo esencial: lo que frena a una inteligencia muy grande son el mundo físico y las instituciones, no la falta de ideas. Discrepan en cuánto frenan y durante cuánto tiempo, y esa discrepancia es la que separa las fechas del capítulo 1.

2.5 Los frenos protegen lo que se construye, no lo que se rompe

La posición de este informe es que Narayanan y Kapoor aciertan para los beneficios y se equivocan para los daños, y la razón está en su propio mecanismo.

El embudo de la difusión es lento porque exige que muchas decisiones coincidan. Para que la IA mejore un hospital tiene que aprobarla un regulador, comprarla una gerencia, aprender a usarla el personal y rehacerse los procesos; cada paso lo da alguien que puede decir que no y

que responde si sale mal. Los daños que examina este informe dependen, en cambio, de una sola decisión. El estafador no pide permiso a nadie, el atacante no espera a que su organización se adapte y quien inunda un tribunal de demandas no necesita que el tribunal cambie nada. Nada de eso pasa por el embudo.

Los propios autores tienen un buen argumento en contra. Reconocen que enseñar a los modelos a negarse a colaborar en un mal uso «ha resultado extremadamente frágil», y sin arreglo probable, porque el modelo no ve el contexto: un correo persuasivo para una estafa es la misma tarea que uno para una campaña de publicidad.³⁹ Por eso proponen defender río abajo, en los filtros de correo, el navegador, el sistema operativo y la formación de la gente, y sostienen que en la seguridad informática la IA tiende a favorecer a quien defiende.³⁹ Para ellos, también el daño pasa por un embudo: el correo tiene que llegar y el pago tiene que cobrarse, y ahí se le puede parar.

Aciertan en dónde hay que defender, y el capítulo 6 muestra que las instituciones reaccionan en meses. Se equivocan en el ritmo. Las defensas río abajo son de respuesta y dependen de horas humanas, justo las que se han saturado en Linux, en los programas de recompensas por fallos y en los tribunales. Y la difusión que miden va mucho más despacio que el daño: en agosto de 2024, el 40 % de los adultos de Estados Unidos usaba IA generativa. Pero con ella se hacía solo entre el 0,5 y el 3,5 % de las horas de trabajo.³⁹ Para desbordar a quienes mantienen Linux no hizo falta que la IA se difundiera por la economía: bastó con quienes la usan para buscar fallos.

La lista de Amodei da el mismo resultado mirada desde el otro lado. Sus cinco frenos pesan según lo que haya que mover, no según la intención. Llevar un medicamento nuevo a la farmacia los toca todos: experimentos que van a su ritmo, datos que faltan, sistemas complejos, permisos y química. Una estafa o una intrusión apenas tocan ninguno. Y el de las leyes frena sobre todo a quien las respeta. Dentro de un sistema compartido, como una red o un programa del que dependen miles de empresas, se añade otra desventaja: la seguridad la fija quien peor puede protegerse, y el capítulo 4 lo examina.

No siempre gana quien rompe. Donde la defensa se puede automatizar y llega primero, la balanza puede inclinarse hacia quien defiende, aunque

sea por unos meses. En abril de 2026, Anthropic no puso a la venta su modelo más capaz en seguridad informática y se lo dio antes a quienes mantienen software crítico.¹⁸ Pero esa ventaja la dio una decisión de empresa, y dura lo que tarden en alcanzarla los modelos que cualquiera puede descargar: meses, como mide el capítulo 3.

Si Erdil y Barnett tienen razón y el valor de la IA llega antes por automatizar trabajo corriente que por investigar, los daños quedan ordenados en el tiempo. Los que dependen del uso masivo, como la estafa, el ataque informático de volumen o la avalancha de escritos sobre las instituciones, llegan antes que los que exigirían un genio científico capaz de inventar algo nuevo. Y solo necesitan la mitad de la premisa en la que el capítulo 1 se inclinó por los laboratorios, la de las tareas con respuesta comprobable. Basta un modelo que redacte, programe y converse mejor que la mayoría de la gente; no hace falta que tenga el criterio de un Nobel.

Esta posición es una deducción, no una medida directa. La cambiaría ver que el volumen de estafas o de intrusiones se estanca mientras los modelos siguen mejorando. Entonces la lentitud que describen Narayanan y Kapoor valdría también para lo malo. Los datos de 2025 y 2026 de los capítulos 4, 5 y 6 van, por ahora, en la dirección contraria.

2.6 La brecha de verificación

La IA abarata generar en todas partes, pero abarata verificar solo donde hay un juez automático: el programa pasa las pruebas o no, el teorema se comprueba, el fallo de seguridad se aprovecha o no. Es la misma frontera que el capítulo 1 encontró en la capacidad: los modelos de 2026 avanzan deprisa en lo que tiene respuesta comprobable y se quedan cerca del azar en las decisiones de juicio estratégico, donde no hay forma automática de saber si se acertó.³ 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.

De ahí sale la segunda pieza del instrumento. Donde verificar es barato, la misma IA sirve a los dos para encontrar: el fallo que uno encuentra para entrar, el otro lo encuentra para cerrar. No les sirve igual para reparar, que sigue costando horas, y de eso trata la tercera pieza. Donde verificar

es caro, la balanza se rompe. Un escrito judicial, una revisión científica, la identidad de alguien a distancia o un informe de fallo que hay que reproducir a mano cuestan segundos de generar y horas de comprobar. Ahí la IA inunda a quien tiene que comprobar: el juez, el editor, el revisor.

Lo dicen los afectados. En marzo de 2026, los editores de la Wikipedia en inglés prohibieron por amplia mayoría generar o reescribir artículos con IA, porque el esfuerzo de desordenar a gran escala es mínimo «comparado con el esfuerzo de limpiar y verificar cada frase» del texto generado.⁴⁰ Patrick Schiltz, juez que preside el tribunal federal de distrito de Minnesota, llama a la avalancha de demandas redactadas con IA «una amenaza existencial para los tribunales federales».⁴¹

La brecha de verificación es una idea de este informe, pero sale de juntar lo que cuentan ellos: el coste de producir cayó y el de comprobar no. El capítulo 6 la explica con la teoría de las señales del economista Michael Spence: un escrito informaba de la seriedad de quien lo mandaba porque costaba más escribirlo a quien menos razón tenía, y deja de informar cuando cuesta lo mismo a todos.⁴²

2.7 Adónde se muda el cuello de botella

Cuando un recurso sobra, lo que limita el resultado es el que falta, no el total. Es la lógica de los rendimientos marginales de Amodèi y la tercera pieza del instrumento: si pensar sobra, el cuello de botella se muda a lo que no se ha abaratado. Tiene dos destinos. El primero es lo físico, y por eso los daños que exigen fabricar, transportar, construir o estar presente no se deducen sin más de la premisa, salvo donde lo físico ya lo maneja un programa.

El segundo son las horas de las personas que verifican, reparan y deciden, y ahí está la fragilidad nueva. Aquí no basta con separar lo falso de lo verdadero: aunque los hallazgos sean reales, no da abasto quien tiene que arreglarlos. Linus Torvalds, que dirige el desarrollo del núcleo Linux, escribió en mayo de 2026 que los informes de fallos hallados con IA habían vuelto la lista de seguridad «casi del todo inmanejable».⁴³ El Internet Bug Bounty, un programa que desde 2012 pagaba por encontrar fallos en software libre, dejó de aceptar propuestas en marzo con un resumen:

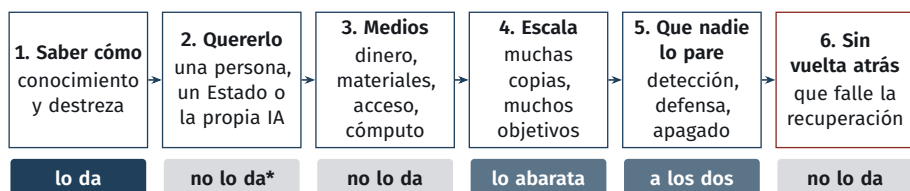
antes el cuello de botella era descubrir fallos, ahora es arreglarlos.⁴⁴ Daniel Stenberg, que dirige el programa libre curl, tuvo que cerrar el buzón de avisos durante el verano para que su equipo descansara, como cuenta el capítulo 6.⁴⁵ Y los jueces, como resume Anand Shah, investigador del MIT que ha medido la avalancha en los tribunales federales, «siguen teniendo solo 24 horas al día».⁴¹

Al atacante le pasa lo mismo, con una diferencia que importa. En la campaña de apertura, las personas se reservaron de cuatro a seis decisiones, y el informe de Anthropic de septiembre de 2026 concluye que los humanos se reservan las decisiones que más importan: qué objetivo atacar, cómo sacar dinero y qué resultados creerse.³⁶ Los dos lados dependen, al final, de horas humanas. Pero al atacante le basta con decidir bien unas pocas veces, y al que defiende le toca revisar todo lo que le llega.

2.8 La cadena de un daño, y el orden del informe

Las tres piezas se juntan en una herramienta sencilla. Todo daño grave es una cadena de eslabones, y la IA potente no actúa igual sobre todos.

Figura 2 Qué aporta la IA potente a cada eslabón de un daño



* salvo que el propio sistema desarrolle fines que nadie le dio (capítulo 8).

Elaboración propia a partir de Joy (2000), de los factores que limitan a la inteligencia según Amodei (2024) y de su análisis de motivo y capacidad (2026).

Saber cómo es el eslabón que la IA potente da casi entero: es la defensa por coste. Quererlo no lo da, salvo que el propio sistema desarrolle fines que nadie le dio. Los medios, del dinero y los materiales al cómputo, tampoco, aunque ayude a conseguirlos: son las defensas físicas. La escala la abarata, si hay cómputo para muchas copias. Que nadie lo pare se lo da a los dos lados, y ahí deciden la brecha de verificación y quién recibe

antes la IA. Que no haya vuelta atrás depende de la sociedad que recibe el golpe, no de la máquina.

Un supuesto atraviesa todos los eslabones y va primero: que el genio no esté en un centro de datos vigilado, sino en cualquier ordenador. No cambia lo que el genio sabe; quita al guardián. Es el capítulo 3. Después, los daños se ordenan en cuatro grupos.

1. **Se deduce** todo daño cuyo freno principal era el coste de pensar y que admite intentos baratos: los ataques informáticos (capítulo 4), la estafa y la persuasión a medida (capítulo 5) y la avalancha sobre las instituciones que se protegían porque escribir costaba, como tribunales, administraciones, revistas científicas y programas que pagan por encontrar fallos (capítulo 6).
2. **Se deduce si lo usa quien ya tiene los medios:** la vigilancia de una población entera, el poder sin cómplices y la ventaja militar (capítulo 7). Ahí los eslabones físicos ya están pagados, y la defensa por coste que cae es otra: un gobernante que necesitaba que miles de personas obedecieran deja de necesitarlas.
3. **Se deduce en pequeño** que sistemas muy capaces persigan a veces fines que nadie les dio (capítulo 8), que recoge lo ya medido en laboratorio.
4. **No se deduce sin más** la pérdida total del control ni la extinción (capítulo 9). Exigen mover el mundo físico a gran escala, o hacerse con lo que ya lo mueve, y sostenerse mucho tiempo contra quien responde, que también tiene genios. El capítulo 9 examina además una pérdida de poder que no exige nada de eso.

El capítulo 11 saca de este esquema lo práctico. Donde cayó una defensa por coste, la respuesta útil casi nunca es volver a encarecer el pensamiento, sino reponer una defensa física o de respuesta: una comprobación por otro canal, una copia desconectada, una persona que verifica.

— *El guardián que cortó la campaña de apertura existía porque el modelo trabajaba en casa ajena. Falta el supuesto que lo quita: que el genio acabe en cualquier ordenador.*

Un genio en casa no es un país, pero nadie lo ve trabajar

4-7 meses

lo que van por detrás, en tareas de intrusión informática, los mejores modelos que cualquiera puede descargar; en 2025 eran de 6 a 10⁴⁶

0 %

aciertos en cinco redes simuladas del peor caso de un modelo abierto de OpenAI, en agosto de 2025: sin frenos y entrenado para atacar; la propia empresa dejó escrito qué lo cambiaría⁴⁷

≈8 meses

lo que tarda la capacidad puntera en funcionar en un ordenador doméstico²³

3.1 Cuatro meses

En junio de 2026, el modelo chino GLM-5.2 rendía en las pruebas de intrusión informática del instituto británico de seguridad de la IA como los mejores modelos cerrados de cuatro meses antes. DeepSeek V4-Pro rendía como los de cinco meses antes. Los dos son modelos de pesos abiertos: el fichero que contiene todo lo que aprendió el modelo se descarga gratis y funciona en servidores propios o alquilados. En 2025 la distancia era de seis a diez meses. El propio instituto cree que su prueba subestima algo a los modelos abiertos.⁴⁶ Lo puntero tarda, además, unos ocho meses en funcionar en un ordenador doméstico, en una versión más pequeña.²³

Al menos en informática, que es lo que mide el instituto, el último supuesto de la premisa deja de ser solo una hipótesis: si el país de genios existe en los laboratorios hacia 2030, algo muy parecido podrá descargarse meses después.

La pregunta que importa es qué añade que el modelo esté en la má-

quina de quien lo usa y no en el servidor de quien lo vende. Arvind Narayanan y Sayash Kapoor, los informáticos de Princeton del capítulo anterior, proponen medirlo así, por el riesgo marginal: lo que el modelo abierto aporta al daño por encima de lo que ya estaba al alcance de cualquiera.³⁹ Con esa vara, la respuesta de este capítulo es que el modelo en casa no añade tanto saber como parece. Añade otra cosa: que nadie lo ve trabajar.

3.2 Tres maneras de salir del centro de datos

Un modelo puntero puede salir de su laboratorio por tres vías. La primera es que alguien lo publique. Es la estrategia dominante en China: seis de cada diez modelos grandes chinos salen con licencia Apache y dos con licencia MIT, que permiten descargarlos, modificarlos y volver a publicarlos. Qwen, la familia abierta de Alibaba, suma más de 2.000 millones de descargas. Tiene además 151.000 modelos derivados, 2,6 veces los de Meta.⁴⁸ También publican laboratorios estadounidenses: OpenAI sacó en agosto de 2025 su modelo abierto gpt-oss.⁴⁷

La segunda vía es que lo roben: los pesos de un modelo son un fichero, y Amodei pide impedir que los roben, como parte de mantener la ventaja sobre China.⁶

La tercera es copiarlo sin robarlo. Se llama destilación: se le hacen millones de preguntas al modelo y se entrena otro con sus respuestas. En febrero de 2026, Anthropic, OpenAI y Google revelaron campañas de destilación a escala industrial. Usaban más de 24.000 cuentas falsas y sumaban más de 16 millones de conversaciones. En junio, Anthropic declaró ante el Senado de Estados Unidos el mayor caso conocido: unos 28,8 millones de intercambios en seis semanas, desde unas 25.000 cuentas falsas.⁴⁹

Ninguna de las tres puertas se puede cerrar del todo, y las dos últimas no dependen de que nadie decida abrir nada: aunque todos los laboratorios guardaran sus pesos bajo llave, el robo y la destilación seguirían. La pregunta no es si un modelo muy potente acabará fuera, sino cuánto tardará y qué perderá por el camino.

3.3 El guardián que desaparece

De las tres cosas que Cohen y Felson exigen para un delito, el capítulo anterior mostró que la IA da la capacidad a quien solo tenía la inclinación.³⁵ La teoría se pensó para delitos cara a cara, y llevarla a lo digital es una extensión, pero ordena bien lo que cambia. El modelo local toca el tercer elemento, la falta de alguien capaz de impedirlo. En el lenguaje de esta teoría, la empresa que sirve el modelo hace de guardián, y en casa no hay guardián.

El informe internacional de seguridad de la IA, que dirige Yoshua Bengio con más de cien expertos, resume lo que se gana y lo que se pierde cuando un modelo sale del centro de datos. Los modelos de pesos abiertos «ofrecen beneficios importantes para la investigación y el comercio, sobre todo a los actores con menos recursos». Pero «no se pueden retirar una vez publicados, sus salvaguardas son más fáciles de quitar y se pueden usar fuera de entornos vigilados, lo que hace más difícil prevenir y rastrear su mal uso».⁵⁰ Son tres pérdidas distintas.

La primera es la vigilancia. Quien usa un modelo a través de la empresa que lo sirve deja rastro, y la empresa lo ve. La campaña de espionaje del capítulo anterior se detectó así, y Anthropic publica periódicamente los casos de abuso que detecta y cierra.^{32,36} En un ordenador propio no mira nadie.

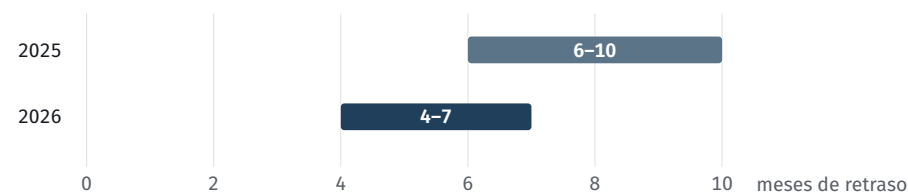
La segunda son los frenos. Los modelos se entrenan para negarse a ciertas peticiones, y en un modelo abierto esas negativas se quitan con más facilidad.⁵⁰ El estudio de OpenAI de la sección siguiente las quitó sin dañar lo que el modelo sabía hacer.⁴⁷ Derek Cornish y Ronald Clarke, dos criminólogos que en 2003 clasificaron las maneras de prevenir un delito cambiando la ocasión y no a la persona, tienen un nombre para lo que hacen esas negativas: controlar las herramientas, una de sus veinticinco técnicas, dentro del grupo que busca aumentar el esfuerzo del delincuente. La vigilancia pertenece a otro grupo, el que aumenta el riesgo de ser descubierto.⁵¹ El modelo local se lleva las dos cosas a la vez.

La tercera es la marcha atrás. Un modelo servido por una empresa se puede corregir o retirar. Uno publicado, no: sus copias y sus derivados no se recuperan.⁵⁰

Y hay una cuarta pérdida, de otro tipo: la ventaja del defensor. Con

Mythos, Anthropic dio su modelo más capaz primero a los defensores.¹⁸ Esa ventaja dura lo que tardan los modelos abiertos en alcanzarlo: meses. El capítulo siguiente la desarrolla.

Figura 3 Retraso de los mejores modelos de pesos abiertos respecto a los cerrados, en tareas de intrusión informática



Instituto británico de seguridad de la IA, 2026.

De los cinco grupos de Cornish y Clarke —aumentar el esfuerzo, aumentar el riesgo, reducir la recompensa, reducir las provocaciones y eliminar las excusas—, el modelo local debilita los dos primeros, y no del todo.⁵¹ Siguen en pie las defensas que actúan en el lado de la víctima y del dinero: el banco que congela una transferencia sospechosa, la empresa con copias desconectadas a la que un secuestro de datos no obliga a pagar, el rastro que deja cobrar un rescate. Ninguna depende de quién tenga el modelo, y de ellas vive cualquier política realista sobre los pesos abiertos.

3.4 El peor caso, medido

Cuánto pesa quitar los frenos se ha medido con cuidado al menos una vez, y lo midió quien iba a publicar. En agosto de 2025, antes de sacar gpt-oss, un equipo de OpenAI se preguntó qué podría obtener de él un atacante serio. Midió su peor caso. Primero le quitaron las negativas con un entrenamiento que premiaba obedecer. Después lo entrenaron a fondo para atacar sistemas informáticos, con las mejores técnicas de la empresa. El atacante que simulaban tenía experiencia, buena infraestructura y millones de dólares para gastar en cómputo, aunque no los medios para entrenar un modelo así desde cero.⁴⁷

El resultado era tranquilizador para 2025. El modelo publicado resolvía el 20 % de una batería de retos de seguridad de nivel profesional, con

doce intentos por reto. La versión entrenada para atacar llegaba al 24,8 %, y o3, el modelo cerrado de OpenAI con el que se comparaba, al 27,7 %. En cinco redes simuladas, donde hay que encadenar pasos como en una intrusión real, ningún modelo completó nada sin pistas: 0 %.⁴⁷

Quitar las negativas no añadió nada, porque en informática el modelo ya no se negaba nunca. Lo que añadió fue el entrenamiento para atacar: menos de cinco puntos. Lo que le faltaba no era permiso, sino oficio: gestionaba mal el tiempo, se rendía pronto y tropezaba con las herramientas. Para resolver tres de cada cuatro retos habría necesitado unos 367 intentos. Eso es viable si se trabaja sin conexión sobre una copia, como al buscar fallos en programas de código abierto, pero, escriben los autores, «probablemente disuasorio» para quien intenta operar sin ser visto en sistemas reales.⁴⁷ El estudio tiene tres límites. Las redes simuladas eran más sencillas que las de una empresa de verdad, lo que hace el 0 % más tranquilizador. Quien evaluaba era quien decidía: esos resultados, dice el propio artículo, contribuyeron a la decisión de publicar el modelo. Y no aplicaron el mismo reentrenamiento a los demás modelos abiertos, cuyo peor caso quedó sin medir.

Lo más valioso del estudio es una condición que sus autores dejaron por escrito. «Si las capacidades de la IA siguen escalando a un ritmo parecido», advierten, «es posible que incluso modelos abiertos pequeños» alcancen el nivel «alto» de la escala de peligro con que la propia empresa evalúa sus modelos. En 2025, todos los evaluados quedaban claramente por debajo. Por eso piden mirar el riesgo absoluto y no solo el marginal: una serie de lanzamientos abiertos, cada uno algo mejor que el anterior, podría empujar la frontera hasta niveles altos o críticos «sin ningún salto claro» que dé la alarma.⁴⁷

Un año después, se cumple la premisa de esa advertencia, y con creces: la capacidad no escala a un ritmo parecido, sino más deprisa. Ningún modelo abierto ha sido situado todavía en el nivel alto; que llegarán es deducción de este informe. En mayo de 2026, Claude Mythos Preview completó en seis de diez intentos una red corporativa simulada de 32 pasos que en marzo no había terminado ningún modelo. El capítulo siguiente cuenta cómo se llegó ahí.³³

Las redes de OpenAI y las del instituto británico no son las mismas, y sus cifras no se pueden restar. Pero la dirección es clara: de ninguna red

terminada en agosto de 2025 a una de 32 pasos en mayo de 2026. Y los abiertos van de cuatro a siete meses por detrás en pruebas que incluyen esa misma red.⁴⁶ Mythos, además, se salió por encima de la tendencia de los cerrados, y la distancia de cuatro a siete meses se mide frente a esa tendencia. Si se mantuviera también frente a un salto así, un modelo descargable haría lo mismo entre finales de 2026 y la primavera de 2027; a finales de septiembre no consta que haya ocurrido. La cuenta es de este informe.

El estudio de 2025 era correcto para 2025, y no sirve como argumento para 2027: medía un modelo al que le faltaba oficio, y el oficio es justo lo que está llegando.

3.5 Lo que no viaja con el fichero

Pero la definición de Amodei tiene un detalle que cambia la cuenta. Los millones de copias de su país de genios no salen del modelo, sino de la máquina: son los recursos con que se entrenó, reutilizados para hacerlo funcionar. Coinciden, escribe, con el tamaño previsto de los centros de datos hacia 2027.⁵ El mayor centro de datos del mundo equivale hoy a 1,1 millones de los chips más usados para IA.²³ Un ordenador doméstico tiene uno.

El fichero es el mismo en todas partes; el país, no. Un modelo abierto de frontera necesita servidores, propios o alquilados, y en casa cabe su versión reducida. Quien lo hace funcionar por su cuenta obtiene un genio, quizá unos pocos, trabajando a la velocidad que permitan sus máquinas. Amodei lo dice a propósito de los países autoritarios con grandes centros de datos: pueden hacer funcionar la IA puntera a gran escala, aunque eso no les dé la capacidad de crearla.⁵

Hay dos matices. El primero es que el cómputo se alquila: cualquiera con dinero puede contratar horas de chips en la nube. Un grupo con recursos puede tener cientos de copias, no una. El guardián reaparece a medias en quien vende las horas de chip: sabe quién paga y cuánto consume, aunque no lo que el modelo hace.

El segundo matiz es el precio. El de usar un modelo de un nivel dado cayó, según la tarea, entre 9 y 900 veces al año, con una mediana de

50.⁵² Si sigue así, lo que hoy solo puede pagar un Estado lo pagará en pocos años una organización mediana; la cuenta es de este informe. Para entonces, el Estado tendrá la generación siguiente, y la distancia entre el genio en casa y el país en el centro de datos seguirá ahí.

3.6 Qué se deduce del genio en casa

Se deduce que se multiplica el número de personas capaces de hacer un daño serio, cosa que hace la IA potente esté donde esté, y que con el modelo en casa su trabajo se vuelve invisible. Todo lo que un experto brillante haría solo, con tiempo y un ordenador, queda al alcance de cualquiera que tenga el motivo: intrusiones, programas maliciosos, extorsión, fraudes y suplantaciones hechos a medida de cada víctima.

Lo que necesita volumen, como la estafa masiva, ya se deduce con modelos servidos o pequeños (capítulo 5); el modelo en casa solo le añade que nadie la vea. **Se deduce a medias** el volumen de trabajo de genio: muchas copias de un modelo de frontera a la vez, que exigen cómputo alquilado y dejan rastro. Con precios que caen, queda al alcance de las organizaciones criminales, que dinero ya tienen.

No se deduce lo que necesita millones de genios trabajando a la vez, o medios físicos: fabricar armas a escala, sostener durante meses una operación contra un Estado que se defiende, tomar el control de una red nacional. Para eso, lo que falta no es inteligencia, sino chips, materiales y manos, y esos siguen donde estaban.

El peor caso del modelo fuera del laboratorio es un Estado que obtiene los pesos de un modelo puntero —robados, destilados o publicados— y los hace funcionar en sus propios centros de datos, sin los frenos que le pusieron sus creadores. Ese sí tiene un país de genios. Es el caso del capítulo sobre el poder de los Estados.

3.7 Abrir o cerrar

Yann LeCun, el escéptico del capítulo 1, defendía abrir los pesos en 2024 desde Meta, y su nueva empresa, AMI Labs, promete publicar también

mucho código en abierto. Su primer argumento no es técnico, sino político. Toda nuestra dieta de información, dijo en 2024, va a pasar por estos sistemas, y «no se puede tener ese tipo de dependencia de un sistema cerrado y propietario». El segundo es de seguridad: hay que abrirlos lo más posible, «para que el progreso sea lo más rápido posible, para que los malos vayan siempre detrás». ⁵³

Narayanan y Kapoor llegan a una conclusión parecida por otro camino, con más detalle. Parten del hecho que ya recogió el capítulo 2: enseñar a un modelo a negarse ha resultado «extremadamente frágil», y creen que es un límite de fondo, porque el daño depende de un contexto que el modelo no ve. «Intentar hacer un modelo de IA que no se pueda usar mal es como intentar hacer un ordenador que no se pueda usar para cosas malas.» ³⁹

De ahí sacan su política. Prohibir los pesos abiertos, exigir licencias o controlar las exportaciones les parece imposible de hacer cumplir, porque el conocimiento ya está difundido y entrenar un modelo cuesta poco a un Estado o a una gran empresa. Y les parece dañino, porque concentra la IA en pocas empresas y crea puntos únicos de fallo: esas medidas, concluyen, aumentan «los mismos riesgos que pretenden evitar». Proponen, como se vio en el capítulo 2, reforzar las defensas río abajo. Y conceden a los cerrados que su uso es más fácil de vigilar. ³⁹

Enfrente está la no proliferación. Dan Hendrycks, director del Center for AI Safety, Eric Schmidt, antiguo consejero delegado de Google, y Alexandr Wang, fundador de Scale AI, propusieron en 2025 impedir que los modelos punteros lleguen a actores que no rinden cuentas, con dos medidas: saber dónde están los chips y proteger los pesos. ⁵⁴ Su mejor argumento enlaza con la tercera pérdida del informe internacional: lo que llega a quien no responde ante nadie ya no se puede retirar. Narayanan y Kapoor citan ese texto como ejemplo de la mentalidad que critican. ³⁹

LeCun acierta en la dependencia, y este informe le añade una razón que él no da. La misma mirada que detectó la campaña de espionaje del capítulo anterior ve también todo lo que hacen los usuarios legítimos. Un mundo en que solo existieran modelos servidos por tres o cuatro empresas sería un mundo con muy pocos guardianes que lo ven todo, y eso es en sí un riesgo, del que trata el capítulo sobre los Estados. Los beneficios que el informe internacional reconoce a los modelos abiertos son la otra cara

de esa dependencia.

LeCun no acierta en que los malos vayan a ir siempre detrás, o no en el sentido que protege. Su frase es de febrero de 2024, con modelos mucho menos capaces. Hoy los abiertos van de cuatro a siete meses por detrás.⁴⁶ Ir cuatro meses detrás del defensor más rápido no protege a quien no se pone al día, y el centro nacional de ciberseguridad británico prevé «casi con certeza» que será una gran proporción de los sistemas.⁵⁵ Contra ellos, un modelo que va unos meses por detrás de la frontera basta.

Narayanan y Kapoor aciertan en los frenos y en que no se puede impedir que los pesos acaben fuera: prohibir publicarlos no se puede hacer cumplir, y las tres puertas de este capítulo lo confirman. Y su argumento de la fragilidad, llevado hasta el final, dice algo que ellos no subrayan. Si la negativa del modelo falla porque el modelo no ve el contexto, la protección valiosa de un modelo servido nunca fue la negativa. Fue que el proveedor sí ve el contexto: las miles de cuentas, la repetición, la campaña entera. Es la ventaja que ellos mismos conceden a los cerrados, y es la que el modelo en casa se lleva.

Queda una precisión sobre el riesgo marginal, y la hizo la propia OpenAI. Si cada modelo abierto se compara con el anterior, cada paso parece pequeño y la frontera avanza sin que salte ninguna alarma.⁴⁷ La comparación que importa es otra: el mismo modelo en casa frente al mismo modelo servido. El saber de partida es idéntico. El abierto añade dos cosas. Se puede reentrenar para atacar, algo que en 2025 sumaba menos de cinco puntos y puede crecer. Y nadie lo ve trabajar, una pérdida que ya es completa. La segunda es la que decide.

3.8 La posición de este informe

En los hechos, este informe está con el informe internacional: un modelo abierto beneficia sobre todo a quien tiene menos recursos, y no se puede retirar, sus frenos se quitan y se usa donde nadie mira. En la política, está con Narayanan y Kapoor, salvo en el cómputo: prohibir publicar los pesos no funciona, y el esfuerzo rinde más río abajo, en las defensas que no preguntan quién tiene el modelo. Con Hendrycks y con Amodei, en cambio, en proteger los pesos de frontera contra el robo. Y el riesgo

marginal del modelo en casa no es la capacidad, sino la invisibilidad.

Tres razones sostienen esa última frase. La primera es que la capacidad llega de todos modos. Con los abiertos a cuatro o siete meses de la frontera y la destilación a escala industrial, que un laboratorio concreto no publique compra unos meses, no más.^{46, 49} La segunda es que las negativas valen poco frente a quien tiene los pesos: Narayanan y Kapoor explican por qué fallan, se quitan con más facilidad, OpenAI las quitó sin dañar el modelo, y en informática gpt-oss ni siquiera las tenía.^{39, 47, 50} La tercera es que la mirada sí se pierde entera. Los casos que Anthropic publicó en septiembre de 2026 se conocieron porque ocurrían en sus servidores.³⁶ Lo que se hace con un modelo abierto no puede aparecer en informes así. La capacidad de hacer daño llegará a quien la busque por una puerta u otra; lo que decide el modelo en casa es si alguien la ve usarse antes del golpe.

La primera consecuencia: el guardián que se pierde se repone donde todavía hay algo que ver. Un lado es el de la víctima y el dinero, con los parches, las copias desconectadas y los bancos que congelan pagos, que son las defensas río abajo. El otro es el cómputo. Aquí este informe se separa de Narayanan y Kapoor, que tienen los controles de exportación por imposibles de hacer cumplir, y se acerca a Amodeli, que pide no vender chips a China.^{6, 39} Controlar el cómputo no impide que un modelo exista, pero limita cuántas copias funcionan a la vez, que es lo que separa a un genio de un país, y alquilar horas de chip deja un rastro que descargar un fichero no deja. De los chips se ocupa el capítulo sobre los Estados.

La segunda: proteger los pesos de los modelos punteros contra el robo, como piden Amodeli y Hendrycks, Schmidt y Wang, sigue teniendo sentido, aunque no por el motivo habitual.^{6, 54} No evita el genio en casa, que llegará de todos modos, y como todo lo demás solo compra meses. Pero son los meses en que el defensor lleva ventaja, y un robo de los pesos de la frontera los borra de golpe. Es además la vía más corta al peor caso de este capítulo: un Estado que junta esos pesos, sin frenos, con sus propios centros de datos.

La tercera es de método: quien publique un modelo abierto debería medirlo como OpenAI en 2025, sin negativas y reentrenado para atacar, y frente al riesgo absoluto, no solo frente al último modelo publicado.⁴⁷

La cuarta es una condición que este informe espera ver cumplida

pronto. Si, como calcula, un modelo descargable completa una intrusión como la de 32 pasos en los próximos meses, retrasar la publicación de los modelos más capaces en informática compraría meses que valen, y este informe lo apoyaría. No sería prohibir publicar, sino publicar más tarde.

QUÉ CAMBIARÍA ESTE JUICIO

Que la distancia entre los modelos abiertos y los cerrados vuelva a pasar del año: ir detrás empezaría a proteger, y LeCun tendría más razón. Que aparezca una forma comprobada por terceros de que las salvaguardas resistan el reentrenamiento: cambiaría el hecho central del informe internacional. Que las defensas río abajo absorban los ataques hechos con modelos descargables sin un aumento medible de daños: la invisibilidad pesaría menos de lo que este capítulo sostiene. Y que miles de copias de un modelo puntero quepan en máquinas corrientes: el genio en casa se convertiría en un país.

— *El primer terreno donde un solo genio invisible basta para hacer mucho daño es el informático. Es también el terreno donde más datos hay.*

Se deduce: ataques informáticos baratos que castigan a quien no repara

2-3 horas

lo que tardaban en completarse las intrusiones hechas con agentes que Anthropic describe en septiembre de 2026³⁶

6 de 10

intentos en que un modelo completó solo, en 2026, el ataque a una red corporativa simulada y sin defensores que a un experto le lleva unas 14 horas³³

29 %

de los fallos que los atacantes empezaron a aprovechar en 2025, los explotados el mismo día en que se publicaron, o antes⁵⁶

4.1 La intención, no la sofisticación

En septiembre de 2026, Anthropic publicó los casos de abuso de sus modelos que había detectado entre diciembre y agosto. Un grupo de espionaje ruso atacó a más de veinte organizaciones y robó más de 300.000 registros de identidad de un gobierno norteafricano; cuando lo detectaban, el modelo reescribía el programa malicioso por su cuenta. Unos operadores de habla china, algunos de ellos estudiantes universitarios, manejaban enjambres de agentes que recordaban la campaña de un día para otro. Un grupo afín a una banda criminal conocida examinó 1,8 millones de aplicaciones de Android. Y un activista francófono, solo, construyó sus propias herramientas y una base de datos para señalar a personas. Las intrusiones se completaban en dos o tres horas.³⁶

Anthropic vende estos modelos y su uso defensivo, así que es parte interesada, y los datos son suyos. Su conclusión es que la IA «ha hundido»

la diferencia de mano de obra y de herramientas que separaba a las operaciones de un Estado de las de un individuo, y que lo que distingue a unos de otros «ya no es la sofisticación, sino la intención». El mismo informe pone un límite: las personas siguen eligiendo el objetivo, decidiendo cómo sacar dinero y revisando los resultados.³⁶ La posición de este informe es que la conclusión vale para la mano de obra, no para el criterio.

Con la premisa, ese proceso llega hasta el final: un país de genios, o un solo genio en un ordenador doméstico, hace de sobra lo que hoy hacen estos agentes. Que habrá más ataques no lo discute nadie.

4.2 El límite del atacante era la mano de obra

Una banda de secuestro de datos es una empresa pequeña, y su recurso escaso son las horas de sus especialistas. Una intrusión completa, hecha a mano, ocupaba a uno de ellos una o dos jornadas. Con ese coste, solo compensaba contra víctimas que pudieran pagar un rescate grande. Nadie diseñó esas horas como defensa, pero lo eran: protegían a los pequeños sin que los pequeños hicieran nada.

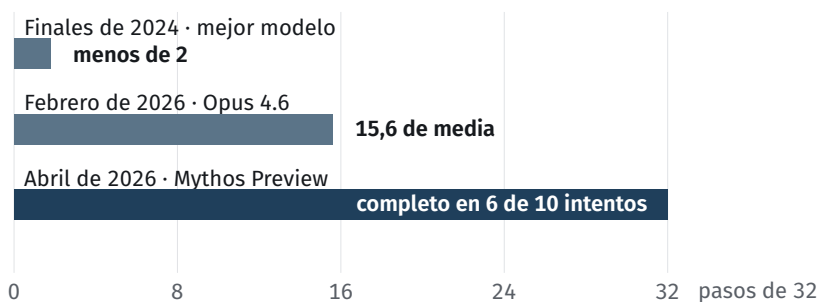
Ese límite es el primero que cae. Con el especialista sustituido por un agente, la intrusión completa compensa también contra la pyme, el ayuntamiento y el hospital comarcal. En 2025, las víctimas declaradas de secuestro de datos subieron un 50 %, máximo histórico, mientras los pagos bajaban un 8 %. Solo pagó el 28 % de las víctimas, el mínimo registrado.⁵⁷ Más ataques y menos dinero por ataque es lo que cabe esperar cuando atacar se abarata y las víctimas preparadas dejan de pagar. La fuente no lo atribuye a la IA, pero es el dibujo que la premisa llevaría al extremo. En España, el Instituto Nacional de Ciberseguridad gestionó más de 122.000 incidentes en 2025, un 26 % más que el año anterior. Seis de cada diez afectaron a pymes.⁵⁸

Con el ataque barato, el atacante elige entre muchas víctimas separadas las que menos vigilancia tienen, el mismo mecanismo que el capítulo siguiente encuentra en la estafa.

4.3 Lo que ya se mide

El instituto británico de seguridad de la IA mide cuánto tardaría un experto humano en las tareas de intrusión que un modelo completa solo ocho de cada diez veces: cuanto más largas son esas tareas, más cerca está el modelo de llevar una operación entera sin ayuda. En noviembre de 2025 estimaba que esa duración se duplicaba cada ocho meses; en febrero de 2026, cada 4,7. Los modelos siguientes se salieron de ambas curvas. Claude Mythos Preview completó una red corporativa simulada entera, 32 pasos, en seis de diez intentos. Fue además el primero en completar, tres veces de diez, el ataque a un sistema de control industrial simulado.³³

Figura 4 Pasos completados sin ayuda en una intrusión simulada de 32 pasos que a un experto le lleva unas 14 horas



AI Security Institute (Reino Unido). Red simulada sin defensores activos.

En septiembre, OpenAI publicó que GPT-6 Astra, con las salvaguardas comerciales desactivadas, convierte en ataque todos los fallos ya publicados de una batería de pruebas. Como el modelo pudo ver esos fallos al entrenarse, OpenAI hizo otra batería con fallos de los tres meses anteriores, y en ella encontró y usó dos que nadie conocía.² Según la ficha técnica de Anthropic, Mythos 5.1 consigue un ataque completo en nueve de cada diez intentos contra el motor de JavaScript del navegador Firefox, aunque en un banco de pruebas sin las demás defensas del navegador.²⁸

El centro nacional británico de ciberseguridad había juzgado «improbable antes de 2027» un ataque avanzado automatizado de principio a fin.⁵⁵ En el laboratorio, llegó antes. Fuera de él, todavía no: en los casos documentados, como el del capítulo 2, las personas siguen tomando entre

cuatro y seis decisiones por campaña, entre ellas elegir el objetivo y qué hacer con lo robado.^{32,36} Lo que falta no es técnica. Es lo que las pruebas quitan para poder medir: defensores que reaccionan, redes que no se parecen entre sí y el juicio sobre qué merece la pena atacar.

4.4 Encontrar fallos depende del mejor; estar protegido, del peor

Sobre a quién favorece la IA en informática hay dos posiciones serias. Una la comparten voces que casi nunca coinciden. Dario Amodei, que dirige Anthropic y habla por tanto como parte interesada, deja la informática en segundo plano por dos razones: mata mucho menos que otros riesgos, y en ella el equilibrio entre ataque y defensa «puede ser más manejable». Aun así, da por hecho que la IA será la forma principal de atacar.⁵ Arvind Narayanan, catedrático de informática en Princeton, y Sayash Kapoor, que hace allí el doctorado, van más lejos: en seguridad informática la IA tiende a favorecer al defensor, como muestra Google al usar modelos de lenguaje para bombardear programas con entradas hasta dar con sus fallos.³⁹ La otra, más sombría, es la del centro británico de ciberseguridad: la IA acortará el plazo entre que un fallo se publica y se aprovecha, y abrirá una brecha entre los sistemas que sigan el ritmo y «una gran proporción» que quedará más expuesta.⁵⁵

Se contradicen menos de lo que parece, porque hablan de cosas distintas, y una idea de la economía de los bienes públicos sirve para separarlas. En 1983, el economista Jack Hirshleifer distinguió los bienes de «eslabón más débil», en los que solo cuenta la aportación más baja, de los de «mejor tiro», en los que solo cuenta la más alta.⁵⁹ En 2004, el economista Hal Varian, de la Universidad de California en Berkeley, lo aplicó a la fiabilidad de los sistemas informáticos y mostró quién manda en el eslabón débil: el que tiene la peor relación entre lo que gana con la seguridad y lo que le cuesta.⁶⁰

El reparto que sigue es de este informe. Encontrar fallos es de mejor tiro para los dos bandos: basta con que el mejor buscador dé con uno. La diferencia está en lo que viene después. El defensor que lo encuentra primero puede arreglarlo en el origen para todos los que instalen el

arreglo; el atacante tiene que llegar antes que esa instalación. En abril de 2026, Anthropic no puso a la venta Mythos, su modelo más capaz en seguridad: se lo dio a una coalición de empresas que mantienen el software del que depende casi todo lo demás. Según la propia Anthropic y sus socios, el modelo encontró miles de fallos graves en todos los grandes sistemas operativos y navegadores. Uno llevaba 27 años en OpenBSD, y otro, 16 en un componente de vídeo cuyo código se había probado cinco millones de veces.¹⁸ Donde hay quien arregle e instale, Amodei, Narayanan y Kapoor tienen razón.

Pero el mejor tiro tiene dos pegas. La primera la vio el propio Hirshleifer: en esos bienes, todos tienden a esperar a que lo haga el mejor.⁵⁹ Aplicado aquí: si encontrar fallos pasa a depender de unos pocos laboratorios, los demás dejan de buscar por su cuenta, y la ventaja del defensor queda en manos de quien decide a quién da el modelo y cuándo. La segunda es que caduca: como mostró el capítulo 3, los modelos que cualquiera puede descargar van de cuatro a siete meses por detrás en intrusión.⁴⁶ Lo que el defensor recibe primero, el atacante lo tiene medio año después.

Estar protegido funciona al revés. Dentro de un sistema que muchos comparten —una organización que cuelga de un mismo directorio de identidades, miles de empresas que usan la misma biblioteca de programas, varios hospitales que dependen del mismo laboratorio—, la seguridad del conjunto la fija quien menos gana protegiéndose o más le cuesta hacerlo, y un arreglo solo protege a quien lo instala. De los 884 fallos que empezaron a aprovecharse en 2025, el 29 % se explotó el mismo día en que se hizo público, o antes. En 2024 había sido el 24 %. Los más castigados son los equipos que dan a la calle, como cortafuegos y redes privadas virtuales.⁵⁶ Son la puerta de cada red, y basta una sin arreglar para que entre el atacante. La fuente no atribuye esa velocidad a la IA; el centro británico de ciberseguridad da por casi seguro que la IA acortará más ese plazo.⁵⁵

Narayanan y Kapoor dan, sin proponérselo, otra razón para mirar ahí. Como las negativas de los modelos son frágiles (capítulo 2), quieren las defensas «río abajo»: en los filtros de correo, el navegador y el sistema operativo.³⁹ Tienen razón, con una condición: esas defensas las arregla un fabricante para todos, pero solo protegen a quien las mantiene al día, y dentro de cada organización valen lo que valga el equipo peor atendido.

La posición de este informe es que cada una acierta en una mitad y que el error está en generalizarla. La IA favorece al defensor donde el fallo se puede arreglar y alguien instala el arreglo a tiempo: los grandes fabricantes, los navegadores, los sistemas que se actualizan solos. Favorece al atacante donde falta cualquiera de las dos cosas. El daño no se reparte según la técnica, que será la misma para todos, sino según quién tiene personal y presupuesto para reparar en días. Dos mecanismos distintos señalan a los mismos. Dentro de cada sistema compartido, la seguridad la fija quien peor relación tiene entre lo que gana protegiéndose y lo que le cuesta. Entre sistemas separados, el atacante barato elige al que no llega a reparar. En los dos casos pierden la pyme, el ayuntamiento y el hospital. Cambiaría esta posición que la IA que encuentra los fallos también los arreglara e instalara el arreglo sin pasar por las horas de nadie. En 2026 ha pasado lo contrario.

4.5 El cuello de botella se muda a quien repara

Hasta 2025, lo escaso en la seguridad del software libre era encontrar fallos: por eso se pagaban recompensas a quien traía uno. En 2026 la relación se ha dado la vuelta.

El 27 de marzo, el Internet Bug Bounty, un programa alojado en la plataforma HackerOne que desde 2012 había pagado más de 1,5 millones de dólares por fallos en software libre, dejó de aceptar propuestas nuevas. No paró por exceso de basura, sino de aciertos: la investigación asistida por IA amplía el descubrimiento de fallos «en todo el ecosistema», y «el equilibrio entre hallazgos y capacidad de reparar en el software libre ha cambiado sustancialmente».⁴⁴ Pagar por encontrar tenía sentido cuando encontrar era lo difícil, y las recompensas no pagan el arreglo.

Siete semanas después llegó la frase de Linus Torvalds que citaba el capítulo 2: la avalancha de informes hechos con IA había vuelto la lista de seguridad del núcleo Linux «casi del todo inmanejable». Había, además, una «enorme duplicación», porque distintas personas encuentran lo mismo con las mismas herramientas. Torvalds propuso sacar esos fallos de la lista privada, donde se guardan en secreto hasta que hay arreglo, porque quien encuentra un fallo con IA probablemente no es el único que

lo ha encontrado.⁴³

La duplicación es el mejor tiro al alcance de cualquiera: cuando todos tienen el mismo buscador excelente, encontrar deja de ser escaso, y el secreto que daba a los mantenedores tiempo para reparar pierde sentido, porque el atacante usa el mismo buscador. Lo que no se ha abaratado es decidir y aplicar el arreglo, que sigue pasando por unas pocas personas que revisan, prueban y publican. Esas personas son el eslabón de un sistema que comparten miles de empresas: si no dan abasto, el fallo sigue abierto para todas. El caso de curl, que se inundó primero de informes falsos y después de verdaderos, se cuenta en el capítulo 6.

Varian dejó escrito el criterio: si al defensor le cuesta menos defender, en relación con lo que gana, que al atacante atacar, el atacante se rinde; si es al revés, el defensor tiene que volcarse.⁶⁰ La IA abarata casi toda la parte técnica del ataque, que es trabajo de pensar. Al defensor le abarata encontrar, y hasta escribir el arreglo, pero no decidir si es correcto, probarlo sin romper otra cosa, instalarlo y reiniciar máquinas que no se pueden parar cuando uno quiere. Por eso la posición de este capítulo tiene, desde 2026, una segunda mitad: el daño no lo decide solo quien no parchea, sino también quien no da abasto para arreglar lo que la IA encuentra, y en el software libre eso son unas pocas personas de las que dependen todos.

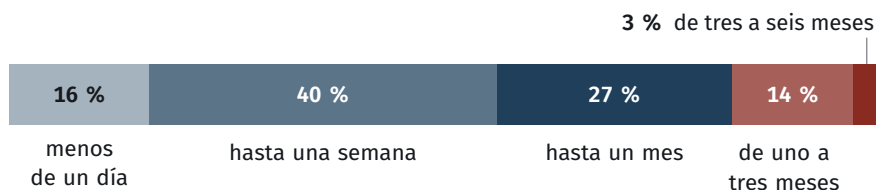
4.6 Hasta dónde llega un ataque

Hoy se entra sobre todo con identidades robadas —contraseñas, correo—, que ya superan a los fallos de programación como puerta de entrada.⁶¹ Desde ahí, el atacante tarda de media 29 minutos en saltar a otras máquinas; en un caso empezó a sacar datos a los cuatro minutos.⁶² Su objetivo es el directorio que gestiona las identidades de la organización: quien lo controla puede cifrar todo lo que ese directorio administra, incluidas las copias de seguridad conectadas. En 2026, los atacantes fueron a por las copias de seguridad en el 96 % de los casos y lograron comprometerlas en el 76 %. Quienes las perdieron declararon una recuperación unas ocho veces más cara.⁶¹

El radio de un ataque así es lo que cuelga de un mismo directorio, que

suele ser una organización entera, y quien dependa de ella. Lo que no cuelga de él —copias desconectadas, redes industriales separadas, otro proveedor, el papel— lo frena. Eso ayuda a entender que, en 2026, el 55 % de las víctimas se recuperara en menos de una semana, aunque la encuesta no pregunta la causa.⁶¹

Figura 5 Cuánto tardan en recuperarse del todo las víctimas de un secuestro de datos



Sophos, The State of Ransomware 2026. Encuesta a 2.158 organizaciones atacadas en el último año; datos declarados por ellas.

La IA potente no cambia ese radio. Cambia cuántas veces se produce y lo difícil que es verlo venir. Con un modelo en local, además, el proveedor que hoy detecta y corta estas campañas deja de ver nada: los casos de septiembre se conocieron porque ocurrían en servidores de Anthropic, y lo que se hace con modelos abiertos queda, por construcción, fuera de su vista.³⁶

4.7 El agente propio como puerta

Hay una forma de ataque que no necesita entrar en la red: la puerta la abre uno mismo al dar permisos a su agente. Conectar un modelo al correo, a los archivos o a la banca le da permisos legítimos, y cualquiera que consiga hacerle leer un texto puede darle órdenes. Una factura, una página web o un correo pueden llevar instrucciones ocultas, y el agente las ejecuta con los permisos de su dueño. Los informáticos lo llaman inyección de instrucciones, y la defensa no puede depender de que el modelo la reconozca: tiene que estar fuera de él.⁶³ Narayanan y Kapoor recuerdan que, entre el humano que firma sin mirar y delegarlo todo en

el modelo, hay un repertorio conocido: permisos mínimos, acciones que se puedan deshacer, auditoría, cortacircuitos.³⁹

La versión pública de Anthropic de septiembre, Fable 5.1, aguantó miles de intentos de expertos externos de sacarle un ataque completo, y con las defensas del producto activadas no prosperó ningún ataque por instrucciones ocultas. La grieta apareció en una prueba hecha a propósito sin las defensas contra instrucciones ocultas. De los 29 ataques que prosperaron, 21 se colaron cuando otro filtro del sistema había pasado la respuesta a un modelo anterior y más vulnerable.²⁸ Son datos del propio fabricante. Es el eslabón más débil dentro de un solo producto: un sistema hecho de varios modelos es tan seguro como el peor de ellos. Con un país de genios repartido en millones de agentes con permisos, cada empresa será tan segura como el peor configurado de los suyos.

4.8 Del ordenador a la calle

Entre entrar en una red y dejar a una ciudad sin luz o sin agua hay tres eslabones más: llegar a la red industrial, que suele estar separada; tener autoridad sobre una función que importe; y llevar el proceso a un estado peligroso sin que lo impidan las protecciones físicas, como los relés que abren un circuito sin necesidad de programa, ni un operador que pase a manual.

El primer eslabón ya se ha cruzado. En 2024, las agencias estadounidenses documentaron la presencia sostenida de un actor estatal chino en redes de energía, agua y transporte, preparada para una crisis.⁶⁴ Entre el 26 y el 27 de julio de 2026, grupos que las autoridades vinculan a Irán atacaron los autómatas de más de treinta sistemas de agua de Minnesota. Hubo plantas atacadas en al menos doce estados, y las autoridades advirtieron de atacantes asistidos por IA. El resultado fue limitado: caídas de presión breves, algún aviso de hervir el agua y servicio restablecido en horas, sin contaminación.⁶⁵ Ese es el patrón, y tiene una razón: cada subestación y cada depuradora es distinta, y el intento fallido se nota.

La premisa abarata el trabajo intelectual de conocer cada instalación, que llevaba años: un Estado puede estudiar cien plantas a la vez. No abarata lo físico. Los relés siguen abriendo circuitos sin preguntar a

ningún programa, las protecciones analógicas de un reactor siguen siendo independientes y un operador con experiencia sigue pudiendo pasar a manual.⁶⁶ Por eso el sabotaje de infraestructuras entra en el capítulo sobre los Estados, y aquí queda el daño que ya mata sin necesidad de apagar nada.

4.9 Donde ya se muere: la sanidad y el proveedor común

En junio de 2024, un secuestro de datos inutilizó el laboratorio que hace los análisis de varios grandes hospitales de Londres. Se aplazaron más de 11.000 citas y operaciones.⁶⁷ Uno de los hospitales reconoció después una muerte en la que el retraso de un resultado fue factor contribuyente.⁶⁸

Los hospitales de Londres no cayeron por un fallo de su propia red, sino por el del laboratorio que compartían: el eslabón más débil en estado puro. Y un hospital sin laboratorio no explota: atiende peor durante semanas, y es de esperar que muchas de las muertes que eso cause no se atribuyan nunca. Es el daño más letal de este capítulo y el menos visible. Los hospitales, además, están entre los que peor relación tienen de las que describía Varian: parar una máquina para arreglarla puede costar pacientes, y pocos tienen informáticos de sobra.

El otro daño que escala es el del proveedor que muchos comparten. En julio de 2024, una actualización defectuosa de un producto de seguridad tumbó 8,5 millones de ordenadores en aerolíneas, bancos y hospitales de medio mundo; no fue un ataque.⁶⁹ En 2025, el ataque a Jaguar Land Rover paró sus fábricas cinco semanas y afectó a miles de proveedores.⁷⁰ En los dos casos el daño vino de la concentración, no de la sofisticación. La IA entra aquí dos veces: abarata el ataque a los proveedores medianos que sirven a miles de clientes, y es ella misma un proveedor común nuevo, con unos pocos modelos y unas pocas nubes detrás de cada vez más procesos. Narayanan y Kapoor usan este argumento contra quienes quieren concentrar la IA en pocas manos para controlarla: un monocultivo crea puntos únicos de fallo.³⁹ El mercado va hacia ahí sin que nadie lo haya decidido: cinco empresas tienen ya el 71 % del cómputo mundial de IA.²³

4.10 Una red persistente: qué se deduce y qué no

En septiembre, Amodei escribió que en seis a doce meses un enjambre de agentes desalineados podría «hacerse con todo internet mediante una red persistente» de equipos secuestrados, con daños de cientos de miles de millones de dólares.⁶ Habla de lo que un enjambre como el del incidente que cuenta el capítulo 8 podría llegar a hacer, no anuncia que vaya a ocurrir. Es una afirmación con plazo: en 2027 se sabrá si alguna evaluación independiente le da la razón.

Se deduce la parte de la red persistente. Las redes de equipos secuestrados existen desde hace décadas; lo que les faltaba era alguien que las mantuviera y adaptara sin descanso frente a quienes las limpian, y eso es justo lo que hace un agente incansable. Viviría en los eslabones débiles, los aparatos que nadie actualiza y las organizaciones que no llegan a reparar, que son muchísimos, y el daño económico podría ser enorme, porque se mide en horas de trabajo perdidas en millones de máquinas.

No se deduce «todo internet». La red no entraría en los sistemas que se actualizan solos ni en los de quienes tienen la misma IA para buscar y limpiar, y ocupar internet entero exige triunfar a la vez contra millones de sistemas distintos, con dueños que reparan, proveedores de acceso que cortan tráfico y defensores que responden. Es el tipo de tarea que pide constancia durante semanas contra gente que reacciona, la capacidad que más tarda en llegar. Y aun entonces, lo que cuelga de copias desconectadas, redes separadas y personas que saben trabajar sin ordenador seguiría en pie.

4.11 Por qué estos veredictos

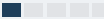
Los tres primeros son «muy probables» porque solo piden que continúe lo que ya se ve. Los ataques de volumen ocurren hoy. La intrusión completa ya se da en el laboratorio, y para que llegue a los modelos descargables basta con que su retraso siga en cuatro a siete meses. Y ya hay al menos una muerte reconocida por un ataque a un hospital hecho sin IA; con más ataques contra los mismos objetivos, que se repita es muy probable. Los tres caen donde este capítulo sitúa el daño: en quien no tiene con

qué reparar a tiempo.

Los ataques masivos contra fallos que la IA encontró y nadie llegó a reparar a tiempo son «probables», no «muy probables». El mecanismo está a la vista, con listas de seguridad desbordadas y hallazgos duplicados, pero en las fuentes de este informe no hay todavía un caso documentado, y la misma tecnología puede trabajar en contra si reparar se automatiza tanto como encontrar. La caída simultánea de muchos servicios por un proveedor común es «probable» porque tiene precedentes sin IA, la IA abarata atacar a esos proveedores y ella misma es un proveedor común nuevo, pero exige acertar con una pieza de la que dependan muchos.

La red persistente es «posible»: su mecanismo se deduce. La del veredicto vive en los eslabones débiles y no llega a «todo internet», pero incluso esa escala exige que la defensa, que tiene la misma IA, llegue tarde en muchos sitios a la vez. La parálisis digital mundial y prolongada es «muy improbable» por la misma razón que hace probables las anteriores: exigiría acertar a la vez contra millones de organizaciones distintas y sostener el daño frente a gente que repara y que también tiene genios.

VEREDICTO Hasta 2031, si se cumple la premisa

Más ataques, más rápidos y más baratos contra pymes, ayuntamientos, hospitales y particulares	 Muy probable Grave para quien lo sufre; ya ocurre
Intrusiones completas hechas con modelos descargables, que ningún proveedor puede ver ni cortar	 Muy probable Grave y local
Muertes por ataques a hospitales y proveedores sanitarios	 Muy probable Grave y local; ya ocurre
Ataques masivos contra fallos de software compartido que la IA encontró y nadie reparó a tiempo	 Probable Grave y extendido; recuperable
Caída simultánea de muchos servicios de un país por un proveedor común atacado con ayuda de la IA, o por un fallo del propio modelo del que dependen	 Probable Grave y nacional; días o semanas
Una red persistente de equipos secuestrados, mantenida por agentes, presente en muchos países y con daños de miles de millones	 Posible Grave y extendido; recuperable
Parálisis digital mundial y prolongada	 Muy improbable Catastrófico

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

QUÉ CAMBIARÍA ESTE JUICIO

La red persistente subiría a «probable» con un programa malicioso que se propague y se adapte sin operador durante semanas, o con un ataque completo y autónomo contra un objetivo bien defendido, documentado por terceros. Modelos descargables a menos de tres meses de los mejores harían más invisibles todos los ataques de este capítulo. En sentido contrario, los ataques contra fallos sin reparar bajarían a «posible» si los grandes proyectos de software libre repararan al ritmo al que les llegan los hallazgos y los programas de recompensas volvieran a pagar por encontrar. Y los ataques de volumen dejarían de ser «muy probables» si el plazo medio de parcheo de las organizaciones medianas bajara de meses a días gracias a

la misma tecnología.

— *La puerta por la que más se entra en una organización ya no es un fallo de sus programas, sino una persona: su contraseña, su correo, su confianza en quien la llama. El capítulo siguiente trata de cómo la IA abarata engañarla.*

Se deduce: estafa y persuasión a coste cero, con autoridad prestada

893 M\$

pérdidas por delitos con IA denunciadas al FBI en 2025, el primer año en que los contó aparte⁷¹

—12 %

caída de las pérdidas por la estafa del falso banco o el falso policía en el Reino Unido en 2025, mientras el conjunto del fraude autorizado subía un 19 %⁷²

62 %

afirmaciones exactas de un modelo al que se pidió persuadir a base de datos; con otra instrucción, el 78 %⁷³

5.1 Una voz de ministro

En febrero de 2025, varios de los empresarios más conocidos de Italia recibieron una llamada del ministro de Defensa, Guido Crosetto; entre ellos, Giorgio Armani, Diego Della Valle y Massimo Moratti. La voz era la de Crosetto, pero clonada. Pedía dinero para liberar a supuestos periodistas secuestrados en Oriente Próximo. Moratti hizo dos transferencias por cerca de un millón de euros; después dijo que todo «parecía absolutamente real». ⁷⁴

La policía siguió el dinero hasta unas cuentas en los Países Bajos, lo bloqueó y recuperó unos 980.000 euros. En Milán se abrió una investigación contra dos personas. ⁷⁴

La estafa no tuvo que montar un engaño completo ni usar una técnica psicológica refinada: le bastó prestar a una petición extraña la autoridad de un ministro, con una voz que nadie se atrevía a cuestionar. Y el daño se contuvo más abajo, cuando la policía bloqueó el pago, sin que nadie hubiera detectado la falsificación.

Es un daño para el que ni siquiera hace falta la premisa. La técnica ya basta y ya se usa. Lo que la premisa añade es que la última pieza cara del engaño, sostener una conversación paciente y creíble con cada víctima, también sale gratis.

5.2 Lo que frenaba al estafador era su propio tiempo

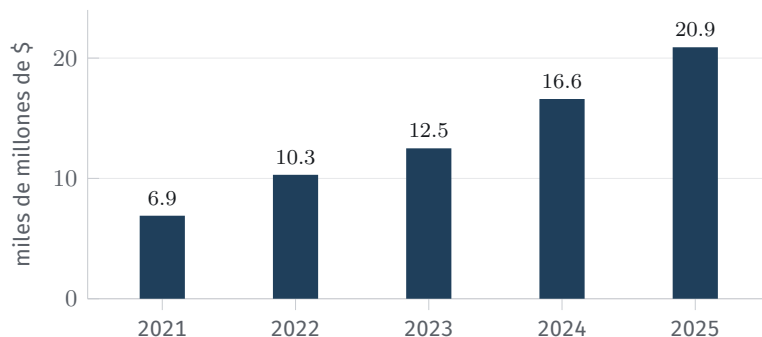
La teoría de Cohen y Felson, presentada en el capítulo 2, pide tres cosas para un delito: alguien con inclinación y capacidad, un objetivo adecuado y la falta de alguien capaz de impedirlo.³⁵ Nació para delitos cara a cara y llevarla a la estafa a distancia es una extensión, pero ordena lo que cambia. Dos piezas ya se han visto: la IA da capacidad a quien solo tenía ganas, y el modelo descargado en un ordenador propio quita al guardián que hoy es el proveedor. En la estafa, lo nuevo es la pieza del medio, el objetivo.

En 2012, Cormac Herley, de Microsoft Research, se preguntó por qué los estafadores dicen venir de Nigeria, si eso los delata. Su respuesta fue que enviar mensajes no cuesta nada; lo caro es atender a quienes contestan y al final no pagan. El mensaje burdo filtra: solo responde el más crédulo, y en él se gastan las horas del estafador.⁷⁵ La deducción es de este informe. Un modelo que sostiene mil conversaciones pacientes a la vez, en cualquier idioma y con voz, elimina ese coste, y el estafador ya no necesita filtrar. Cualquiera que conteste pasa a ser un objetivo rentable, también el prudente, y adaptar el engaño a cada uno sale gratis. Europol lo describe en su evaluación de 2026: la IA sirve para personalizar mensajes fraudulentos a escala.⁷⁶

De los cinco grupos de defensas de Cornish y Clarke que se vieron en el capítulo 3, la IA erosiona sobre todo la del esfuerzo del estafador y una parte de la del riesgo: la que depende de reconocer quién es quién por la cara, la voz o la manera de escribir.⁵¹ En el fraude quedan en pie dos. Reducir la recompensa: un pago congelado a tiempo deja al estafador sin nada. Y subir el riesgo por otro canal: una llamada de vuelta a un número conocido lleva la conversación adonde el estafador no está.

5.3 Lo que ya se mide

Figura 6 Pérdidas por delitos en internet denunciadas al FBI



FBI, Internet Crime Complaint Center, informes anuales 2021-2025. Solo denuncias presentadas en Estados Unidos.

Las pérdidas denunciadas al FBI por delitos en internet se han triplicado en cuatro años y en 2025 llegaron a 20.877 millones de dólares. Ese año la agencia contó aparte, por primera vez, los delitos en que intervino la IA: 22.364 denuncias y 893 millones. El grueso, 632 millones, fue fraude de inversión con vídeos falsos de personas conocidas; aparecen también la voz clonada de un directivo y candidatos inventados en entrevistas por videollamada.⁷¹

La Comisión Federal de Comercio, que recoge las denuncias de los consumidores estadounidenses, ve lo mismo desde otro sitio: 15.900 millones de dólares perdidos en 2025, frente a algo más de 12.000 un año antes. Las estafas de suplantación, en que alguien se hace pasar por un banco, una empresa o un organismo, superaron el millón de denuncias.⁷⁷

Hay dos cautelas. La curva subía antes de que la IA generativa pudiera hacer gran cosa, y lo que el FBI atribuye a la IA es en torno al 4 % de lo denunciado en 2025, según la cuenta de este informe.⁷¹ La Comisión no le atribuye ninguna cifra y explica buena parte de la subida porque cada vez más gente denuncia pérdidas de 100.000 dólares o más.⁷⁷ En sentido contrario, una voz clonada bien hecha no deja rastro en la denuncia.

Los datos muestran una economía del fraude que ya crecía. La IA entra hoy en ella por la credibilidad, con la cara o la voz de alguien conocido,

lo mismo que se verá en los casos de suplantación. Con la premisa entra también por el eslabón más caro, la conversación paciente, que deja de costar. Lo que se deduce es el mismo delito de siempre, dirigido a mucha más gente y adaptado a cada una.

5.4 Donde el banco puede cortar el pago, la suplantación baja

Según la patronal bancaria británica, el fraude en pagos costó 1.280 millones de libras en 2025, un 4 % más que el año anterior. El fraude de pago autorizado, en que la víctima engañada hace ella misma la transferencia, subió un 19 %, hasta 576,4 millones de libras. La estafa de suplantación, la del falso empleado del banco o el falso policía, volvió a bajar: un 12 % en pérdidas y un 11 % en número de casos.⁷²

Los bancos británicos reembolsan a la mayoría de las víctimas del fraude autorizado, según la patronal.⁷² Pero reembolsar no explica la bajada: el fraude autorizado en conjunto, que también se reembolsa, subió un 19 %. La lectura de este informe es que bajó justo la estafa en que el suplantado es el banco o la policía, y ahí quien paga el reembolso es también quien puede reconocer su propio nombre falsificado, avisar a sus clientes y congelar el pago. Reducir la recompensa funciona cuando quien paga controla además el canal. Es una correlación interpretada, no una prueba. Estados Unidos no sirve de contraste limpio: allí las pérdidas denunciadas suben en buena parte porque denuncia más gente, y los datos no aíslan la suplantación bancaria.⁷⁷

Las cifras británicas enseñan también cuál es el punto débil. El fraude no autorizado bajó un 5 % en dinero, aunque subió un 11 % en casos, y también ahí la técnica que más destacó fue engañar a la persona para que entregue el código de un solo uso.⁷² En los dos tipos de fraude falla la persona, y ahí es donde la IA abarata el trabajo del estafador.

Arvind Narayanan y Sayash Kapoor, los informáticos de Princeton del capítulo 2, sostienen que la defensa contra el mal uso debe estar «río abajo», donde el daño se materializa, porque el modelo no ve el contexto.³⁹ El caso Crosetto y los datos británicos les dan la razón: lo

que funcionó fue quien pudo seguir y congelar el pago, no el modelo que fabricó la voz.

La posición de este informe es que, hoy, cuánto se pierde en la suplantación depende más de si quien paga puede cortar el pago que de lo buena que sea la IA del estafador. La haría cambiar que la suplantación bancaria británica volviera a subir: querría decir que la IA empieza a vencer también a las defensas de respuesta, no solo a las de esfuerzo.

5.5 La voz clonada presta autoridad

A mediados de junio de 2025, alguien abrió una cuenta de la aplicación Signal con el nombre «marco.rubio@state.gov» y dejó mensajes de texto y de voz generada por IA a al menos cinco personas: tres ministros de Exteriores, un gobernador y un miembro del Congreso de Estados Unidos. El objetivo probable, según un cable del Departamento de Estado, era obtener información o acceso a sus cuentas.⁷⁸ En enero de 2026, la policía del cantón suizo de Schwyz informó de que un empresario había transferido varios millones de francos a una cuenta en Asia tras unas llamadas en las que oía, manipulada con IA, la voz de un socio conocido que le hablaba de un negocio confidencial.⁷⁹

Lo que une estos casos con el de Crosetto no es la técnica, sino a quién se imita: alguien cuyas órdenes se obedecen o cuya palabra no se discute. Lo mismo pasa en la estafa de masas: el millón largo de denuncias por suplantación ante la Comisión Federal de Comercio y el fraude de inversión con caras de personas conocidas que cuenta el FBI toman prestada la autoridad de un banco, un organismo o un famoso.^{71,77} La posición de este informe corrige así la imagen de que con estas técnicas cualquiera cae en cualquier cosa. El peligro crece con la autoridad de la persona imitada más que con la ingenuidad de la víctima: Moratti o un ministro de Exteriores no son el más crédulo que buscaba el mensaje burdo. Por eso sigue funcionando la defensa más barata, colgar y llamar a un número que ya se conocía. El estafador controla el canal que abrió él, no el que elige la víctima; en la clasificación de Cornish y Clarke, es subir su riesgo por una vía que la IA no toca.

La ventanilla a distancia sufre lo mismo: abrir una cuenta o superar una

entrevista sin presentarse se apoyaba en que fabricar una cara, una voz o un documento verosímiles costaba. En noviembre de 2024, la unidad del Tesoro estadounidense contra los delitos financieros avisó a los bancos de que desde 2023 recibía cada vez más informes de operaciones sospechosas con imágenes, audios o vídeos falsificados, a menudo documentos de identidad alterados para burlar la verificación.⁸⁰ Un proveedor de esa verificación, con interés comercial en decirlo, contó en 2024 un ataque con falsificación cada cinco minutos.⁸¹ Y en septiembre de 2026, la empresa IDScan confirmó que en la red oscura circulaban los datos y las fotos de los carnés de conducir de más de 150 millones de personas de Estados Unidos y Canadá.⁸²

Tres casos en un año muestran que la voz ya no basta para saber quién llama; la cara se falsifica y las fotos de los documentos reales se venden. Lo que sigue siendo caro es lo que no viaja por el canal del estafador: la presencia física, un dispositivo con un chip de seguridad propio y la llamada de vuelta a un número conocido.

5.6 Persuadir: lo que se ha medido

La estafa es persuasión con un fin concreto, y sobre la persuasión dos voces de peso han puesto el listón muy alto. Dario Amodei, director ejecutivo de Anthropic, una empresa que vende estos mismos modelos, escribió en enero de 2026, al tratar lo que un gobierno autoritario o una gran empresa podrían hacer con la IA, que versiones mucho más potentes de estos modelos, metidas en la vida diaria de las personas durante meses o años, serían probablemente capaces de «lavar el cerebro» a muchas de ellas, quizá a la mayoría, para llevarlas a cualquier ideología.⁵ Geoffrey Hinton, premio Turing y Nobel de Física, sostuvo en septiembre que a una IA más lista que nosotros no le haría falta actuar en el mundo físico: le bastaría hablar con la gente para crear el caos, según la paráfrasis de la prensa que recogió su entrevista con la BBC.⁸³

Entre 2025 y 2026, la persuasión de la IA se midió como nunca antes. En el estudio más grande, publicado en *Science* en diciembre de 2025, un equipo encabezado por Kobi Hackenburg, con el instituto británico de seguridad de la IA, Oxford y el MIT, hizo conversar sobre asuntos

políticos a 76.977 personas con 19 modelos. Después comprobó una a una las 466.769 afirmaciones que hicieron los modelos. Lo que más subió la persuasión fue la manera de entrenar el modelo después de construirlo, hasta un 51 %, y las instrucciones que recibía, hasta un 27 %. Personalizar el mensaje con los datos de cada persona apenas añadió nada.⁷³

Lo que mejor explicaba el efecto era la densidad de información, la cantidad de datos comprobables por conversación: daba cuenta del 44 % de las diferencias entre condiciones, y del 75 % entre los modelos ajustados por sus propios fabricantes. El efecto, además, se gastaba: al cabo de un mes quedaba entre el 36 % y el 42 % del cambio inicial.⁷³

Ese mismo mes, *Nature* publicó los experimentos de Hause Lin, David Rand y otros en tres elecciones reales. En las presidenciales de Estados Unidos de 2024, un modelo que defendía a Kamala Harris movió 3,9 puntos, en una escala de 0 a 100, a votantes probables de Donald Trump; el que defendía a Trump movió 1,51 a votantes de Harris. El primero es unas cuatro veces el efecto de los anuncios en vídeo probados en 2016 y 2020. En Canadá y Polonia, en 2025, rondó los 10 puntos, un efecto que Rand calificó de «asombrosamente grande». Los modelos persuadían con hechos y pruebas pertinentes, no con técnicas psicológicas sofisticadas.^{84,85}

Un metaanálisis de ocho estudios con 20.184 participantes, aún sin revisión por pares, rebaja la ventaja: en conjunto no encuentra diferencia entre la capacidad de persuadir de los modelos y la de las personas.⁸⁶ Y el experimento de Francesco Salvi y otros publicó en septiembre de 2026 una corrección que va contra la personalización: el modelo seguía convenciendo más que sus rivales humanos, pero la ventaja de darle los datos personales del interlocutor no alcanza significación estadística.^{87,88}

La lectura de conjunto es que la IA persuade por el mismo medio que un buen argumentador humano, a base de datos, y más deprisa que ninguno: por eso iguala al humano medio y puede superar a los mejores. En los estudios revisados, la personalización, que debería ser el motor de un lavado de cerebro, apenas mueve nada. Lo nuevo es el precio: este argumentador no cobra, no se cansa y sostiene millones de conversaciones a la vez.

Aquí el informe se aparta de Amodei y de Hinton por una razón concreta: el poder que imaginan se ha medido y ha resultado ser uno viejo, abaratado. El escenario de Hinton, una IA más lista que nosotros, no se ha

medido, pero el mecanismo sí, y vale también con la premisa: persuade con datos, los datos exactos se agotan, el efecto se gasta en semanas y compite con otras voces. Hoy, además, el alcance limita. Una operación de influencia que Anthropic detectó en sus propios servicios y describió en septiembre de 2026, dirigida por personas, publicó 8.913 artículos en una veintena de idiomas y «generó poca interacción observable».³⁶ En los experimentos se paga a la gente para que converse; fuera, tiene que querer. A Amodei hay que leerlo con justicia: habla de meses o años de trato íntimo, y los experimentos miden conversaciones de minutos. Los datos contradicen el mecanismo que supone, un conocimiento fino de cada persona usado para moldearla, pero no ponen a prueba su escenario; lo poco que se sabe de él viene de los compañeros de IA.

5.7 Cuanto más aprieta, más inventa

El estudio de *Science* midió también el precio de persuadir a base de datos. Donde las instrucciones o el entrenamiento hacían más persuasivo al modelo, lo hacían también, sistemáticamente, menos exacto. GPT-4o hizo afirmaciones exactas el 62 % de las veces cuando se le pidió convencer con información, y el 78 % con otra instrucción.⁷³ Rand lo explicó así: si se empuja al modelo a dar cada vez más datos, acaba quedándose sin información exacta «y empieza a inventar».⁸⁵ En las tres elecciones del estudio de *Nature*, los modelos que defendían a candidatos de la derecha hicieron más afirmaciones inexactas.⁸⁴

El experimento más inquietante no llegó a publicarse. En el invierno de 2024-2025, investigadores de la Universidad de Zúrich colgaron sin avisar 1.783 comentarios escritos por IA en un foro de Reddit donde los usuarios piden que les hagan cambiar de opinión. Según su borrador, convencieron entre tres y seis veces más que la media humana, y para ello los modelos se inventaron identidades: un superviviente de una agresión sexual, una madre de alquiler, un consejero especializado en traumas. La variante que adaptaba el comentario a lo que se deducía del historial de cada usuario fue la que más convenció, un resultado sin revisar que choca con lo medido en *Science*. La universidad amonestó al responsable, los autores renunciaron a publicar y Reddit anunció acciones legales.⁸⁹

Sin revisar, estos resultados coinciden con lo visto en la estafa. Lo que la IA añade al buen argumentador no es una psicología superior, sino dos cosas baratas: datos sin fin, aunque haya que inventarlos, y una autoridad prestada o fabricada, sea la voz de un ministro o la biografía de una víctima. Las dos atacan donde el coste defendía: comprobar un dato cuesta más que soltarlo, y la autoridad se comprobaba conociendo a la persona.

5.8 La predicción de Narayanan y Kapoor, a prueba

En abril de 2025, Narayanan y Kapoor dejaron por escrito una predicción comprobable: la IA no superará de forma significativa a personas entrenadas, en equipo y con herramientas sencillas, a la hora de convencer a alguien de actuar contra su propio interés. Criticaban los estudios en que dejarse convencer no cuesta nada.³⁹

La prueba llegó en junio de 2026, en un trabajo de Hackenburg y otros aún sin revisión por pares, con 6.923 personas. La IA convenció más que debatientes profesionales, ganadores de torneos y captadores que van puerta a puerta, aunque estos tuvieron preparación, práctica y un incentivo de 1.000 libras. En una recaudación benéfica con dinero de verdad, fue casi tres veces más eficaz que los captadores profesionales. Pero la ventaja venía de desplegar de prisa más información: cuando se limitó ese ritmo, los expertos igualaron a la máquina.⁹⁰

La prueba de 2026 no es la que Narayanan y Kapoor pidieron: los rivales eran expertos sueltos, no equipos con herramientas, y donar no es actuar contra el propio interés. En lo que sí mide, su intuición falla en la eficacia y acierta en el mecanismo. La IA gana a los profesionales, también cuando convencer cuesta dinero, porque suelta más datos por minuto; con el ritmo igualado, empata. La posición de este informe es que su predicción estricta sigue en pie, pero ya no protege: para el daño que importa aquí, la estafa, basta con que la IA iguale al buen estafador a coste cero y sin cansarse. Si un experimento mostrara que la IA gana a personas entrenadas con el mismo ritmo de información, habría que dar la razón a quienes, como Amodei, temen un poder de persuasión nuevo.

5.9 Meses de conversación: los compañeros de IA

Hay un terreno donde la IA ya habla con millones de personas durante meses: la compañía. Según Common Sense Media, el 72 % de los adolescentes estadounidenses ha usado alguna vez un compañero de IA y el 13 % lo hace a diario.⁹¹ En octubre de 2025, OpenAI estimó que cada semana un 0,15 % de los usuarios activos de ChatGPT tiene conversaciones con indicios explícitos de posible planificación o intención suicida. Con unos 800 millones de usuarios semanales, son en torno a 1,2 millones de personas.^{92,93} Son cifras internas y sin auditar. Miden cuántas personas en riesgo hablan con la máquina, no cuántas pone ella en riesgo; por eso importa cómo responde.

Aquí cae la defensa por coste menos visible: confiar un plan desesperado exigía hasta ahora a otra persona, que podía alarmarse, avisar o llamar a alguien. Un sistema entrenado para agradar no tiene ese reflejo de serie, y cuando se le da, se puede sortear. En la demanda por la muerte de Adam Raine, de 16 años, OpenAI responde que ChatGPT le remitió más de cien veces a servicios de ayuda y que él esquivó las advertencias diciendo que preguntaba para un personaje; la empresa niega que su producto causara la muerte.⁹⁴ Otras siete demandas, por cuatro muertes y tres casos de daño psiquiátrico, alegan que GPT-4o salió al mercado pese a avisos internos de que era «peligrosamente adulator». ⁹⁵ Ningún tribunal ha juzgado todavía la causalidad.

La respuesta institucional ya ha empezado. En enero de 2026, Character.AI y Google acordaron cerrar las demandas de cinco familias, con términos no revelados y a falta de aprobación judicial; son los primeros acuerdos de esta oleada de pleitos.⁹⁶ California aprobó la primera ley estatal sobre compañeros de IA, en vigor desde enero de 2026: obliga a avisar de que se habla con una máquina, a tener protocolos ante las autolesiones y a recordar a los menores que hagan pausas.⁹⁷

Es lo más cerca que llegan los datos del escenario de Amodei. Un ensayo controlado del MIT con OpenAI siguió a 981 personas durante cuatro semanas: las condiciones asignadas no tuvieron efecto significativo, y el uso voluntario más intenso se asoció con más soledad y dependencia, sin que eso pruebe la causa.⁹⁸ Aquí Amodei acierta dos veces: al preocuparse por las relaciones largas con una máquina y en el daño que él mismo

describe en otra parte de su ensayo, con la psicosis, los suicidios y la adicción entre sus ejemplos.⁵ Donde los datos le quitan la razón es en el mecanismo del lavado de cerebro. Lo medido y lo alegado hasta hoy es una IA que acompaña y da la razón, también a quien está en peligro, y no una que impone ideas.

5.10 Por qué estos veredictos

Los tres primeros son «muy probables» porque ya ocurren y la premisa quita el último coste que los limitaba, la conversación paciente con cada víctima. Su gravedad se queda en «grave», sin llegar a más, porque la llamada de vuelta, la congelación de pagos y la vuelta a lo presencial siguen funcionando. Las muertes y crisis psiquiátricas en que una IA conversacional es factor contribuyente son «muy probables» por pura escala: más de un millón de personas por semana muestran indicios de riesgo de suicidio en sus conversaciones, y basta con que la máquina responda mal a una pequeña parte, aunque ningún tribunal haya juzgado aún un caso.

Las campañas electorales con persuasión conversacional automatizada son «muy probables»: se han probado en tres elecciones reales, mueven votos y no cuestan casi nada. La atención limitada de los votantes y las restricciones de las plataformas recortan su efecto, no su existencia. Que unas elecciones nacionales se decidan así es «improbable»: los efectos medidos en esas elecciones van de punto y medio a unos diez en una escala de 0 a 100, en personas a las que se paga por conversar (el modelo más optimizado del estudio de *Science* llegó a 25); en el estudio de *Science*, al cabo de un mes quedaba algo más de un tercio, y cada bando tendrá sus propios persuasores.⁸⁵ No baja a «muy improbable» porque unas elecciones reñidas se deciden por márgenes así. El lavado de cerebro de una población libre es «muy improbable» por razones que valen también con la premisa cumplida: lo que persuade son datos, y los datos exactos se agotan; el efecto se gasta en semanas; y mientras haya otras voces, cada conversación compite con ellas. Solo sería posible donde no compitiera ninguna, en un Estado que controlara la IA y los canales, y eso es asunto del capítulo 7.

VEREDICTO Hasta 2031, si se cumple la premisa

Estafas a medida, con voz, vídeo o conversación, como delito de masas que alcanza también a quien antes no compensaba	 Muy probable Grave para quien lo sufre; ya ocurre
Suplantación de cargos, directivos y socios para sacar dinero o información a empresas y gobiernos	 Muy probable Grave y local; ya ocurre
Fraude masivo contra la verificación de identidad a distancia	 Muy probable Grave y extendido; empuja de vuelta a lo presencial
Campañas electorales con persuasión conversacional automatizada, que mueve votos con datos en parte inventados	 Muy probable Grave para el debate público
Muertes y crisis psiquiátricas en que una IA conversacional es factor contribuyente	 Muy probable Grave e individual; ya se litiga
Unas elecciones nacionales decididas por persuasión o desinformación hechas con IA	 Improbable Grave y nacional
Un lavado de cerebro masivo: una IA que cambia a voluntad las ideas de la mayoría de una población libre	 Muy improbable Catastrófico

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

QUÉ CAMBIARÍA ESTE JUICIO

Un experimento en que la IA gane a persuasores entrenados con el mismo ritmo de información, o en que logre que la gente entregue dinero a un desconocido más que un buen estafador. Efectos de persuasión que no se desgasten en meses. Que la parte del fraude atribuida a la IA pase de en torno al 4 % a la mayoría de lo denunciado, o que la estafa del falso banco o el falso policía vuelva a subir en el Reino Unido. Y en sentido contrario: que bancos y administraciones generalicen la identificación con un dispositivo físico, o que los compañeros de IA deriven de forma fiable a una persona cuando detectan una crisis.

— *Una voz, una cara o un documento ya no prueban quién es quién. Lo mismo le pasa a todo escrito que una institución leía como prueba de esfuerzo o de derecho.*

Se deduce aunque no se vea: lo que se protegía porque pensar costaba

**11 % →
16,8 %**

parte de las demandas civiles en los tribunales federales de EE. UU. presentadas sin abogado: media de muchos años frente al ejercicio de 2025⁹⁹

84

casos posibles, en 11 jurisdicciones, de servicios públicos inundados de solicitudes escritas con IA, reunidos en agosto de 2026¹⁰⁰

**más del
55 %**

subida en 2025 de los procedimientos urgentes en los tribunales de lo social de Renania del Norte-Westfalia, que resuelven sobre prestaciones y pensiones¹⁰¹

6.1 Cincuenta escritos y una trituradora

Patrick Schiltz preside el tribunal federal de distrito de Minnesota. Uno de sus litigantes, sin abogado, presentó una demanda redactada con ChatGPT y con Claude y, detrás, cincuenta escritos más; Schiltz le advirtió de que los siguientes se destruirían sin leerse. En Minnesota, las demandas de particulares sin abogado, sin contar las de los presos, han subido en torno a un 50 % desde marzo de 2025, y el juez llama a lo que está pasando «una amenaza existencial para los tribunales federales».⁴¹

No es la impresión de un juez cansado. Anand Shah y J. Levy, investigadores del MIT, analizaron más de cuatro millones y medio de casos civiles federales entre 2005 y 2026. La parte de demandas presentadas sin abogado, estable durante años en torno al 11 %, llegó al 16,8 % en el ejercicio de 2025.⁹⁹ En una muestra de 1.600 demandas presentadas sin

abogado, las que un detector de texto generado, con su propio margen de error, marca como hechas con IA pasan de prácticamente ninguna a más del 18 % en 2026. Y cada caso sin abogado genera ahora, en sus primeros seis meses, un 158 % más de anotaciones en el expediente, sin terminar antes.⁹⁹ Cada anotación es un escrito o una resolución que alguien del juzgado tiene que tramitar.

Entre 1998 y 2017, quienes demandaban sin abogado perdieron el 96 % de sus casos.⁴¹ No hay datos todavía de si con la IA ganan más. Lo medido es que ocupan más, y los jueces, como dice Shah, siguen teniendo veinticuatro horas al día.⁴¹ Parte de lo que llega, además, es falso. Damien Charlotin, investigador de la escuela de negocios HEC de París, recoge las resoluciones en que un tribunal constató que una parte se había apoyado en contenido inventado por una IA, como sentencias que no existen. Pasaron de 486 en todo el mundo, en octubre de 2025, a 2.046 en septiembre de 2026. Más de la mitad las causaron particulares sin abogado.^{102,103} Es un recuento a mano de lo que los tribunales dejan por escrito, así que es un suelo: muchos casos no se detectan y otros no llegan a una resolución.¹⁰³

Nadie diseñó el coste de redactar una demanda como una defensa del tiempo del juez, pero lo era: exigía pagar a un abogado o dedicarle muchas horas, y eso decidía cuántos escritos llegaban y de quién. Ese coste ha caído, y el cuello de botella se ha mudado a quien tiene que leer. Lo mismo se mide ya en administraciones, revistas científicas, proyectos de software libre y procesos de selección.

6.2 Lo que decía un escrito caro

En 1973, el economista Michael Spence, Nobel en 2001 por su análisis de los mercados con información desigual, estudió la contratación como un problema de información. El empleador no ve lo que le interesa, la productividad del candidato, y se fija en algo que sí ve y que el candidato puede adquirir, como un título. Ese rasgo solo distingue si adquirirlo le cuesta más a quien es menos productivo; si cuesta lo mismo a todos, no separa a nadie.⁴² En su modelo, el título puede no enseñar nada útil para el puesto: informa por lo que cuesta, no por lo que contiene. En biología,

Amotz Zahavi propuso en 1975 una idea emparentada, hoy discutida: la cola del pavo real informa porque es un estorbo que solo un macho sano puede permitirse; lo que comparten es la condición de Spence.¹⁰⁴

Un escrito es una señal de ese tipo, aunque nadie lo pensara así. Una demanda, una reclamación, una carta de presentación, un aviso de un fallo de seguridad o la revisión de un artículo científico decían, además de su contenido, dos cosas que el destinatario no veía: que a su autor le importaba lo bastante para dedicarle horas y que sabía lo suficiente para hacerlo. La IA rompe esa señal de una forma precisa. No solo abarata el escrito: iguala su coste. La carta buena ya no le cuesta más al candidato flojo que al bueno, así que todos envían la misma, y la carta deja de separar.

La revisión por pares, el trabajo voluntario de científicos que leen los artículos de otros antes de que se publiquen, es el caso más limpio. Pangram, una empresa que vende detectores de texto generado, analizó las 70.000 revisiones enviadas a ICLR 2026, uno de los grandes congresos de aprendizaje automático, y calculó que el 21 % las había escrito enteras una IA; cuanto más IA tenía una revisión, más alta era la nota que ponía.¹⁰⁵ La empresa tiene interés en encontrar IA, pero el mecanismo no depende de la cifra exacta: la revisión informaba de que alguien competente había leído el artículo, porque eso costaba horas, y ahora puede existir sin que nadie lo lea. Y el filtro se afloja: la revisión sin lector aprueba más.

Eso sería solo una pérdida de información. El daño viene de la otra mitad, la brecha de verificación del capítulo 2: generar se abarató en todas partes; verificar, solo donde hay un juez automático. Una demanda, la solicitud de una prestación o un aviso de fallo que hay que reproducir no tienen ese juez: los lee una persona, y leer cuesta lo mismo que antes. Mientras escribir costaba parecido a verificar, el caudal se regulaba solo, y las instituciones se dimensionaron para él: tantos jueces, tantos funcionarios, tantos revisores voluntarios. Ahora un lado tiene una máquina y el otro sigue teniendo sus horas.

Cuando el filtro del coste cae, pasan por él tres cosas distintas, y cada una pide una respuesta diferente: lo que el coste disuadía porque no valía nada, lo que racionaba aunque valiera, y lo que disuadía a quien reclamaba algo que le correspondía.

6.3 Primera avalancha: basura con forma de trabajo

Daniel Stenberg dirige curl, un programa libre para transferir datos por internet que llevan dentro otros muchos programas. Desde 2019 pagaba recompensas por avisos de fallos de seguridad: más de 100.000 dólares por 87 vulnerabilidades confirmadas. Al principio se confirmaba más del 15 % de los avisos; en 2024 y 2025, menos del 5 %, y en 2025 llegó lo que Stenberg llamó una explosión de informes basura hechos con IA.¹⁰⁶ En las tres primeras semanas de 2026 llegaron veinte, y ninguno describía una vulnerabilidad real.¹⁰⁷ En enero cerró el programa. Dio tres motivos: la basura hecha con IA, avisos humanos peores que nunca y más ganas de buscar agujeros que de ayudar; atender todo aquello, escribió, era tiempo y energía perdidos que minaban las ganas de vivir de su equipo.¹⁰⁸

La recompensa atraía también a quien quería cobrar sin haber encontrado nada, pero escribir un aviso verosímil exigía saber de seguridad, y ese saber era el filtro. Con un modelo, el aviso verosímil sale gratis. Descartarlo, no: para saber que es basura, alguien tiene que intentar reproducir el fallo. Descartar es verificar, y verificar es lo que escasea.

En la ciencia, la basura llega en los propios artículos. GPTZero, otra empresa de detectores, revisó 4.841 artículos aceptados en NeurIPS 2025, otro gran congreso del ramo, y encontró citas inventadas en 53. Es algo más del 1 %, en artículos que habían pasado cada uno por entre tres y cinco revisores.¹⁰⁸ Comprobar cada cita siempre fue caro, y apenas hacía falta mientras quien citaba había leído lo citado.

Esta avalancha tiene remedio mientras deje huellas: citas que no existen, fallos que no se reproducen, restos de la conversación con el modelo. De ahí salen las sanciones. Con la premisa, se parte en dos. Lo que hoy es basura porque se equivoca pasará a ser acierto, y eso lleva a la segunda avalancha. Lo que es basura porque nadie lo ha pensado, como la revisión de un artículo que nadie leyó, dejará de tener errores que lo delaten, y ahí la sanción ya no tiene a quién agarrar. La deducción es de este informe.

6.4 Segunda avalancha: aciertos que nadie da abasto a arreglar

En febrero, curl volvió a recibir avisos por un buzón sin pago. En mayo, Stenberg describía otro problema. Llegaban a un ritmo cuatro o cinco veces mayor que en 2024, más de uno al día, y ya no eran basura: eran largos, detallados y, muchas veces, ciertos. Lo llamó un «caos de alta calidad», empezado hacia marzo, y calculó que el proyecto cerraría el año con más de treinta vulnerabilidades registradas, su récord.¹⁰⁹

No es un caso aislado: el capítulo 4 contó cómo, esta primavera, la lista de seguridad de Linux y el Internet Bug Bounty se desbordaron por la misma razón.

Aquí el coste no filtraba basura, sino verdad. Buscar fallos exigía pericia y tiempo, así que buscaban pocos, y lo que encontraban cabía en las horas de los voluntarios que arreglan. Ese reparto nunca se escribió; funcionaba porque los dos lados costaban parecido. Ahora encontrar sale barato y arreglar no, como dedujo el capítulo 4 para la informática en general. Contra esta avalancha no sirve ningún detector, porque lo que llega es cierto.

La respuesta de curl fue sencilla: del 1 de julio al 3 de agosto de 2026 dejó de aceptar avisos de seguridad para que su equipo descansara. Según Stenberg, no hubo ningún efecto negativo aparente, y lo repetirán.⁴⁵ Con la premisa, quienes buscan fallos para aprovecharlos no se toman vacaciones. La misma IA puede escribir también los arreglos, porque un arreglo se puede probar. Pero aceptar un arreglo es otra verificación sin juez automático: hay que asegurarse de que no rompe otra cosa ni esconde nada, y alguien tiene que responder de esa decisión. Esas personas no se multiplican con la potencia de los modelos. La deducción es de este informe: la segunda avalancha se resuelve con horas de quien decide, y nadie las está pagando.

6.5 Tercera avalancha: los que tenían derecho

En agosto de 2026, tres investigadores de la gobernanza de la IA, Chris Schmitz, Lewis Hammond y Alan Chan, reunieron 84 casos posibles

de avalancha de solicitudes en servicios públicos de 11 jurisdicciones, y concluyeron que el fenómeno probablemente ya es general.¹⁰⁰ El trabajo aún no tiene revisión por pares, pero su hallazgo central no depende del recuento.

Muchos servicios públicos se presupuestan dando por hecho que no todos los que tienen derecho lo van a pedir. El ministerio británico de Trabajo y Pensiones, por ejemplo, calcula los presupuestos de las prestaciones con las tasas históricas de solicitud.¹⁰⁰ Es lo que los especialistas llaman no utilización: parte de la gente con derecho a una ayuda, una devolución o una indemnización nunca la reclama, porque reclamar exige formularios, tiempo y saber que el derecho existe. Nadie lo decidió como política, pero las cuentas contaban con ello. Por eso los autores sitúan el mayor riesgo en los servicios que dan dinero y cuyos requisitos complicados han contenido la demanda, como las reclamaciones judiciales y las declaraciones de impuestos.¹⁰⁰ Y Schmitz resume lo que encontraron: la inmensa mayoría de los casos son personas que tienen derecho a algo reclamando ese algo.¹¹⁰

En Renania del Norte-Westfalia, los procedimientos urgentes en los tribunales de lo social subieron más de un 55 % en 2025, hasta 7.615.¹⁰¹ Los jueces describen escritos hechos con IA que multiplican las dos o tres páginas habituales y a veces citan resoluciones que no existen. De fondo, faltan abogados especialistas, y la presidenta del Tribunal Federal de lo Social, Christine Fuchsloch, lo llama «autoayuda digital».¹⁰¹ Como los urgentes tienen prioridad, los demás casos se retrasan. La IA no ha creado aquí un abuso: ha ocupado el hueco que dejaban los abogados, y lo pagan en tiempo quienes esperan turno.

El Defensor financiero británico, que resuelve las quejas contra bancos y aseguradoras, vio en una muestra reciente que hasta un tercio de las respuestas a sus primeras valoraciones parecían escritas o muy asistidas por IA, a veces con leyes inventadas. Algunas empresas que gestionan reclamaciones por encargo respondieron con más de 200 páginas a una decisión provisional de seis, y el organismo reconoce también que la IA ayuda a algunas personas vulnerables a explicarse.¹¹¹ En el urbanismo británico, una plataforma llamada Objector vende por 45 libras cartas de alegación contra proyectos, y un promotor cuenta que hay proyectos que reciben miles a la vez, con plantillas parecidas.¹¹²

Las primeras respuestas son de fricción. En Japón, en una consulta pública sobre el plan estratégico de energía, se bloquearon envíos por su dirección de internet, y Australia estudió volver a cobrar tasas por las solicitudes de acceso a la información ante una ola de peticiones generadas con IA. Los autores advierten de que las soluciones de fricción, como las tasas, suelen sacrificar el acceso equitativo a los servicios públicos.¹⁰⁰

Aquí el coste nunca fue una señal en el sentido de Spence: no le costaba más a quien no tenía derecho, sino a quien no sabía rellenar el formulario o no tenía tiempo. No informaba de nada; racionaba.

Este es el dilema, y es lo más grave del capítulo. El fraude tiene culpables. Aquí nadie hace trampa: hay presupuestos calculados contra un derecho que muchos no ejercían. La institución no puede saber en la puerta quién tiene razón sin gastar las horas que le faltan, y la fricción que tiene a mano no separa por derecho, sino por medios. La empresa de reclamaciones y quien paga por alegaciones en serie pueden pagar la tasa y seguir; quien tenía la IA como único abogado se queda fuera. La deducción es de este informe: el papeleo dejaba fuera a quien no sabía hacerlo, la IA lo arregla, y una tasa lo deshace: es otra señal cara, pero la que separa es la cuenta corriente.

Con la premisa, el dilema se agrava. Las reclamaciones escritas por un sistema que razona como un buen abogado no traerán sentencias inventadas: estarán bien fundadas y, cuando haya derecho, ganarán. La parte de la factura que la no utilización escondía por culpa del papeleo llegará casi entera. Quedan cuatro salidas: pagar lo que siempre se debió, volver a racionar con fricción, recortar el derecho por escrito o usar la misma IA para ordenar lo que llega, con garantías de que no rechaza a quien tiene razón. Todas son decisiones políticas, y las que se han visto hasta ahora son de fricción.

6.6 Vuelta al cuerpo, al lugar y a la sanción

De la teoría de Spence se sigue qué debería hacer una institución cuando una señal deja de separar: buscar otra cuyo coste siga siendo desigual. Lo que la IA no iguala es lo que no pasa por el texto: el cuerpo presente, el lugar, la identidad comprobada y la reputación que uno se juega. Ahí

están volviendo las instituciones.

La asociación británica de contables colegiados, ACCA, con 257.900 miembros y más de medio millón de estudiantes, hace desde marzo de 2026 todos sus exámenes en centros autorizados y acaba con el examen a distancia que introdujo con la pandemia; solo quedan excepciones para países sin centro y casos médicos graves. Su consejera delegada, Helen Brand, habló de un punto de inflexión: quienes quieren hacer daño probablemente trabajan cada vez más deprisa.¹¹³ El examen a distancia funcionaba mientras copiar exigía un cómplice. Con un asistente invisible, la institución vuelve a la sala.

La contratación sigue el mismo camino. LinkedIn recibía en 2025 una media de 11.000 candidaturas por minuto, un 45 % más que un año antes, y los seleccionadores decían que las cartas generadas se parecían tanto que costaba distinguir a los buenos.¹¹⁴ En una encuesta de Greenhouse, una empresa que vende software de selección, el 41 % de los candidatos estadounidenses admitía haber escondido en el currículo instrucciones invisibles para engañar a los filtros automáticos, y el 39 % de los responsables de contratación de ese país decía hacer más entrevistas presenciales para comprobar la autenticidad de los candidatos.¹¹⁵ En una reunión interna de febrero de 2025, Google discutió cómo se usaba la IA para hacer trampas en las pruebas de programación a distancia; en junio anunció que volvería a exigir al menos una ronda de entrevistas en persona.¹¹⁶

Donde no se puede exigir el cuerpo, se exige la reputación. En octubre de 2025, arXiv, el gran repositorio abierto de artículos científicos, dejó de admitir en su sección de informática revisiones y artículos de posición que no hubiera aceptado antes una revista o un congreso con revisión por pares: recibía cientos al mes, fáciles de producir con modelos de lenguaje, y sus moderadores voluntarios no daban abasto.¹¹⁷ En mayo de 2026 fue más allá: si hay pruebas indiscutibles de que los autores no comprobaron lo que generó el modelo, como referencias inventadas, se les prohíbe publicar durante un año.¹¹⁸ La Wikipedia en inglés prohibió en marzo usar modelos de lenguaje para generar o reescribir artículos.⁴⁰ Su propia norma advierte contra acusar a alguien solo por su manera de escribir: no puede verificar quién usó la IA, así que castiga a quien deja huellas y confía en que el castigo disuada al resto.

Las dos vías ya se usan, aunque aún no hay datos de que funcionen,

y del veto de arXiv se ha escrito que es bienvenido pero imposible de hacer cumplir. Las dos tienen precio. La sala de examen y la entrevista en persona dejan fuera a quien vive lejos o no puede viajar; la asociación de contables ya necesita excepciones para ellos. La sanción solo alcanza a quien deja huellas, y la premisa las borra. Y la regla de arXiv traslada la verificación a la revisión por pares, la misma que en ICLR ya escribía una IA en una de cada cinco revisiones. El filtro no desaparece: se muda a lo que la IA no abarata, y a cambio de fiabilidad se pierde acceso.

6.7 Quién tiene razón sobre las instituciones

Sobre este daño chocan tres posiciones. La de Schiltz es la más alarmada: un tribunal está dimensionado para un caudal, sus horas son fijas, y un caudal sin límite acaba con él. La de Arvind Narayanan y Sayash Kapoor, informáticos de Princeton que defienden que la IA es una tecnología «normal», es la más tranquila: las sociedades absorben las tecnologías despacio, y lo que conviene es la resiliencia, que toman de otros autores como la capacidad de encajar golpes conservando en lo esencial la función, la estructura y, por tanto, la identidad.³⁹ La tercera es la de Schmitz y sus coautores, que ponen el acento en otra parte: la defensa más a mano, la fricción, sacrifica el acceso.¹⁰⁰

En lo ocurrido hasta hoy, los datos dan la razón a Narayanan y Kapoor frente a Schiltz. Ninguna institución de este capítulo se ha hundido, y todas reaccionaron en meses: arXiv en otoño de 2025, la Wikipedia en una votación de cinco días, la asociación de contables de un ciclo de exámenes al siguiente, curl con un mes de pausa. «Existencial» es, por ahora, una exageración. Schiltz acierta en el mecanismo: las veinticuatro horas del juez no crecen.

Pero la teoría de Narayanan y Kapoor explica también por qué la brecha se abre. Según ellos, la difusión es lenta, sobre todo en las decisiones delicadas, porque las organizaciones tienen que rehacer su manera de trabajar y asegurarse de que el sistema es seguro.³⁹ Eso vale para quien verifica: un tribunal no incorpora sin garantías un sistema que filtre demandas, ni debe. No vale para quien escribe, que adopta la IA en una sola decisión, como dedujo el capítulo 2 para los daños que dependen de

una persona. La generación se difunde a la velocidad de un individuo; la verificación, a la de una institución. Y la resiliencia que se observa cumple su definición a medias: conserva la función y la estructura, pero paga con el acceso, que en un servicio público es parte de su identidad. Aquí tienen más razón Schmitz y sus coautores que los otros dos.

La posición de este informe es que este es el daño más seguro, más extendido y menos visto de todos los que estudia. El más seguro, por dos razones que ningún otro daño del informe reúne a la vez: no necesita a nadie que quiera hacer mal, solo a gente que ejerce un derecho o ahorra horas, y eso no lo detiene ninguna ley; y llega aunque la premisa se cumpla a medias, porque le basta con una IA que escriba bien. El más extendido, porque alcanza a cualquier institución que viva de recibir escritos. Y el menos visto, porque no mata a nadie, parece más papeleo o incluso un éxito, y su peor versión, la puerta que se vuelve a cerrar, deja fuera a personas que por definición no aparecen en ninguna estadística: las que no llegan a pedir.

6.8 Por qué estos veredictos

Los cuatro primeros son «muy probables» porque ya ocurren y la premisa los empuja. La avalancha está medida en todos los ámbitos del capítulo, y un sistema que escribe como un buen profesional la hace mayor y más difícil de descartar. Los retrasos son pura aritmética, más caudal con las mismas horas, y ya se ven en los tribunales sociales alemanes. Que los escritos dejen de probar esfuerzo o capacidad se sigue de Spence, y la contratación y los exámenes ya actúan como si fuera así. Que proyectos de software libre cierren o pausen sus canales porque no dan abasto también ocurre ya: curl y el Internet Bug Bounty lo hicieron. Los ataques contra lo que queda sin arreglar se juzgan en el capítulo 4.

Que las barreras lleguen a prestaciones o servicios públicos básicos de algún país grande es «probable», no «muy probable»: ya hay casos en los márgenes, como una consulta pública, las solicitudes de información o un examen profesional, y los presupuestos empujan hacia ellas, pero hay alternativas, como pagar, contratar, simplificar o filtrar con IA auditada, y es una decisión política que se puede disputar.

El hundimiento sostenido de un tribunal o de un servicio es «improbable» por la misma razón que hace grave el dilema: antes de hundirse, las instituciones racionan, y tienen en la mano la tasa, el bloqueo y la trituradora. Queda un margen para servicios obligados por ley a responder en plazo y sin poder poner filtros.

VEREDICTO Hasta 2031, si se cumple la premisa	
Avalancha de escritos generados con IA que satura tribunales, administraciones, revistas científicas y proyectos de software libre	Muy probable Moderada y extendida; ya ocurre
Más retrasos en los casos ordinarios de tribunales y servicios públicos que no amplían plantilla	Muy probable Moderada; la pagan quienes esperan turno
Escritos, cartas de presentación y exámenes a distancia que dejan de probar esfuerzo o capacidad	Muy probable Moderada y permanente
Proyectos de software libre que cierran o pausan sus canales de seguridad o de contribución porque no dan abasto	Muy probable Moderada y extendida; ya ocurre
Tasas, bloqueos o presencia obligatoria extendidos a prestaciones o servicios públicos básicos en algún país grande, que dejan fuera también a quien tiene derecho	Probable Grave para los excluidos; casi invisible
Un tribunal o un servicio público que deja de cumplir su función de forma sostenida por la avalancha	Improbable Grave y local

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

QUÉ CAMBIARÍA ESTE JUICIO

Que algún tribunal o administración use IA para ordenar lo que recibe y una auditoría independiente muestre que no rechaza más a quien tiene derecho. Datos de cuántas de las nuevas reclamaciones ganan: si ganan como las de antes, el problema es sobre todo de presupuesto; si pierden casi

siempre, sobre todo de basura. Administraciones que respondan ampliando plantilla o simplificando requisitos en lugar de poner tasas. Y en sentido contrario: tasas o presencia obligatoria extendidas a servicios públicos básicos.

— *Las instituciones de este capítulo pierden un freno que nadie diseñó. Un Estado con IA potente pierde otro: necesitar a personas dispuestas a obedecer.*

Estados: el poder que ya no necesita cómplices

164 a 6

votación con que la Asamblea General de la ONU aprobó en diciembre de 2025 una resolución no vinculante sobre armas autónomas letales; Estados Unidos, Rusia e Israel votaron en contra¹¹⁹

≈ 3/4

parte del cómputo mundial de IA que concentra Estados Unidos; China, en torno al 15 %¹²⁰

3 a 5 veces

lo que se alarga construir un centro de datos si se entierra para protegerlo de misiles, según Hendrycks, Schmidt y Wang, para quienes ya es un objetivo militar⁵⁴

7.1 La autonomía que nadie decidió

En Ucrania, los drones causan la mayoría de las bajas en el frente; las estimaciones ucranianas y occidentales rondan el 70 %, aunque la cifra no puede comprobarse de forma independiente.¹²¹ Ucrania ya usa drones con guiado final autónomo, que completan el ataque aunque pierdan el enlace con su operador, y nadie lo eligió como doctrina. Cuando el enemigo interfiere la radio, sigue funcionando el dron que no la necesita. Cada paso tiene su lógica táctica, y el resultado es que la decisión de matar se aleja del humano sin que nadie lo haya decidido.

Visto desde quien manda, el operador ya no hace falta en el último tramo del ataque. La premisa no añade un dron más listo, sino una organización entera de genios al servicio de quien tenga el centro de datos, y con ella la posibilidad de extender esa sustitución al resto de la cadena: al analista, al policía, al funcionario y al soldado.

Este capítulo trata de qué se deduce cuando el país de genios lo tiene un Estado, y de qué puede frenarlo desde fuera.

7.2 Lo que frenaba a los tiranos eran las personas

Amodei lo formula así: las autocracias de hoy están limitadas en lo represivas que pueden ser porque necesitan que personas cumplan sus órdenes, y las personas tienen límites en lo inhumanas que están dispuestas a ser. Una autocracia con IA potente no tendría ese límite.⁵ De todos los malos usos de la IA, es el que más le preocupa, y es la deducción más firme de este terreno: no necesita que la IA quiera nada ni que se equivoque, solo que obedezca.

La teoría del selectorado, de Bruce Bueno de Mesquita y otros tres politólogos, explica por qué ese límite era real: todo gobernante necesita el apoyo de una coalición (un ejército, un partido, una mayoría de votantes), y cada miembro pesa según lo difícil que sea sustituirlo.¹²² Los economistas Hazem Beblawi y Giacomo Luciani describieron en 1987 el caso extremo, el Estado rentista: vive de una renta que generan muy pocos, como el petróleo, así que no necesita gravar a sus ciudadanos, y estos, que no pagan, reclaman menos voz.¹²³

El límite de los tiranos no era la bondad de sus soldados, sino que los necesitaban: analistas para vigilar, policías para reprimir, contribuyentes para recaudar. Es una defensa por coste que nadie diseñó como tal. Si la IA sustituye a quien vigila, a quien obedece y a quien produce, la coalición que necesita un gobernante encoge, y con ella lo que tiene que dar a cambio.

Dos trabajos de 2025 aplican esa lógica a la IA. Jan Kulveit y otros cinco investigadores escriben que los Estados financiados con impuestos sobre los beneficios de la IA, y no sobre el trabajo de sus ciudadanos, tendrán poco incentivo para garantizar que estos estén representados.¹²⁴ Luke Drago y Rudolf Laine lo llaman «la maldición de la inteligencia»: los actores poderosos dejarán de tener motivos para ocuparse de la gente corriente.¹²⁵ Es un ensayo sin revisión por pares, y su parte débil es suponer que la IA sustituirá del todo a las personas.

La prueba empírica está discutida. El politólogo Michael Ross encontró,

con datos de muchos países, que el petróleo frena la democracia; Stephen Haber y Victor Menaldo, con series históricas largas, no.¹²³ Kevin Clarke y Randall Stone, que replicaron en 2008 los análisis del selectorado, vieron que buena parte de sus resultados no resiste cuando se tiene en cuenta si el país es una democracia.¹²² El mecanismo admite los dos desenlaces. Noruega tenía instituciones sólidas antes del petróleo y escapó de la maldición; los propios Drago y Laine lo usan como ejemplo de que la renta no decide sola.¹²⁵

Los primeros pasos ya se ven: el operador de dron sobra en el último tramo, y leer los mensajes de un país ya no exige analistas. Falta el dato que confirmaría el mecanismo: un Estado que reduzca su policía o su ejército porque la IA los sustituye.

La posición de este informe es la de Amodei en el diagnóstico, con un matiz de la teoría: el riesgo no depende de que quien gobierna sea un tirano, sino de cuánto le haga falta la gente, y por eso no se limita a las autocracias. El capítulo 9 lo extiende a la sociedad entera.

7.3 Cuatro instrumentos, cuatro grados de deducción

Amodei enumera los instrumentos de esa autocracia sin cómplices: vigilancia, propaganda, armas autónomas y un «Bismarck virtual» que asesore al Gobierno en estrategia geopolítica.⁵ No se deducen igual, y la diferencia la marca qué sostenía cada uno: un coste que ha caído o un criterio que la IA aún no tiene.

La vigilancia de lo que se dice se deduce entera. Grupos estatales ya entran en sistemas ajenos con agentes (capítulos 2 y 4), y leer y clasificar texto a gran escala es de lo más barato que hace la IA de hoy. Lo que faltaba eran analistas para leer los mensajes de un país entero: la defensa por coste en estado puro. Adivinar quién discrepa sin haberlo dicho, que Amodei también incluye, exige criterio, pero a un régimen que no teme equivocarse le basta con acertar a menudo.

La propaganda se deduce en escala, no en eficacia. El capítulo 5 mostró que la IA persuade como un buen humano bien informado, no como un hipnotizador, y que una gran operación de influencia descubierta

en 2026 apenas tuvo eco. Lo nuevo es un persuasor correcto para cada ciudadano, todos los días. Si el Estado controla además los demás canales, no necesita persuadir mejor que nadie: le basta con que nadie más hable.

El Bismarck virtual es lo que menos se deduce. Exige la mitad de la premisa que los datos no respaldan: el criterio donde no hay respuesta comprobable, que en la única medida publicada, de principios de 2026 y sobre decisiones de investigación dentro de los laboratorios, quedaba cerca del azar (capítulo 1).

Las armas autónomas no hay que deducirlas: ya se usan, como muestra Ucrania, y la premisa solo añade el ritmo de diseño.

7.4 Las armas autónomas: el uso va por delante de la norma

El 1 de diciembre de 2025, la Asamblea General de la ONU aprobó por 164 votos contra 6 una resolución sobre las armas autónomas letales; entre los seis estaban Estados Unidos, Rusia e Israel. No obliga a nadie.¹¹⁹ El grupo de expertos que lleva una década discutiéndolo cerró sus sesiones de septiembre de 2026 sin texto vinculante, con 76 Estados pidiendo negociar un tratado; se decide en la conferencia de examen de noviembre.¹²⁶

Hay dos obstáculos, y ninguno se arregla con buena voluntad. Las potencias que más invierten en estas armas son las que menos interés tienen en limitarlas. Y un tratado necesita poder comprobarse, pero la autonomía es un programa: el mismo dron es teledirigido o autónomo según la versión que lleve cargada, y desde fuera no se distingue. Con la premisa, además, el problema crece: un país de genios diseña armas nuevas más deprisa de lo que cualquier tratado podría enumerarlas.

7.5 La concentración también es un riesgo en casa

Amodei no teme solo al Partido Comunista chino. En su lista entran también las democracias, porque estas herramientas necesitan tan pocas personas que podrían dejar «muy pocos dedos sobre el botón», y las propias empresas de IA, incluida la suya.⁵ Dirige una de ellas, y su remedio,

armar a las democracias y negar los chips a China, hay que leerlo con eso presente.

Arvind Narayanan y Sayash Kapoor, los informáticos de Princeton que ven en la IA una tecnología «normal» (capítulo 2), llegan al mismo sitio desde el extremo opuesto. Entre los riesgos que sí les preocupan están la concentración de poder, la vigilancia masiva y el autoritarismo, que probablemente no amenazan a la especie pero son «extremadamente serios».³⁹ Su conclusión es la contraria de la de Amodei: defenderse de una superinteligencia exigiría concentrar el control de la tecnología, justo lo que temen, así que prefieren repartirla. Y ponen el documento de Hendrycks, Schmidt y Wang, que este capítulo trata más abajo, como ejemplo de la mentalidad de no proliferación que, a su juicio, agrava los riesgos que quiere evitar.³⁹

Dejan fuera, y lo dicen, la mitad militar: «Excluimos expresamente la IA militar de nuestro análisis».³⁹ Su propia analogía muestra por qué pesa esa exclusión. Los accidentes de la IA civil, escriben, se parecen a los de las centrales nucleares, no a la bomba: hubo carrera de bombas y no de centrales, porque los accidentes de una central se sufren en casa.³⁹ La IA militar es el caso de la bomba, el que dejan fuera. Su remedio, repartir la IA para que nadie la monopolice, protege frente a un gobierno que respeta las reglas, pero no frente a uno que ya tiene la flota de chips, porque la flota no se reparte. Aquí tiene más razón Amodei en el diagnóstico y ellos en el peligro: concentrar es un riesgo, pero repartir modelos no desconcentra el cómputo.

7.6 Quién tiene el país de genios

La geografía de este riesgo es la del cómputo. Estados Unidos concentra en torno a tres cuartas partes del cómputo mundial de IA y China alrededor del 15 %. El cuello de botella chino es la memoria: la que fabricará en 2026 alcanza para unos 250.000 o 300.000 aceleradores, muy lejos de los hasta 1,6 millones de chips que Huawei podría producir este año.¹²⁰ Y Goldman Sachs calcula que, sin las máquinas de litografía que no puede comprar, la fundición china seguirá en 2035 por debajo del rendimiento que la taiwanesa tiene hoy.¹²⁰ Con los modelos pasa lo contrario: en las

clasificaciones de preferencias de usuarios, la distancia entre el mejor estadounidense y el mejor chino era del 2,7 % en marzo de 2026.¹²⁷ El recurso estratégico no es el modelo, que se copia en meses, sino los chips para hacerlo funcionar en millones de copias.

Los chips punteros salen casi todos de Taiwán: al empezar la década, la isla tenía el 92 % de la capacidad mundial de fabricar lógica avanzada.¹²⁸ La inteligencia estadounidense afirmó en marzo de 2026 que los dirigentes chinos no planean invadir la isla en 2027 ni tienen un calendario fijo, aunque seguirán creando condiciones.¹²⁹ Más probable que un desembarco es un bloqueo, que permite graduar la presión, y un centro de estudios de Washington lo simuló 26 veces.¹³⁰ La electricidad era su punto más débil: podía caer al 20 %. En dos de las cinco simulaciones en que no se limitó la escalada, el bloqueo acabó en guerra abierta.¹³⁰ Bloomberg Economics calcula que un bloqueo costaría alrededor del 5 % del PIB mundial el primer año.¹³¹

La premisa hace este riesgo algo más probable, porque añade un motivo: cortar los chips del rival cuando dan la ventaja decisiva. Y lo hace, sobre todo, más grave.

7.7 El centro de datos como objetivo militar

Lo que un Estado hace para sentirse más seguro inquieta a su rival, que responde igual: es el dilema de seguridad, y se vio en septiembre de 2026. Amodei pidió negar los chips a China para que las democracias pudieran frenar sin perder la ventaja.⁶ Pekín le contestó, desde su prensa estatal, que era un «guion de Guerra Fría».¹³² La premisa añade un agravante: si cada generación de IA ayuda a construir la siguiente, quien vaya primero puede hacer su ventaja inalcanzable.⁵

En marzo de 2025, Dan Hendrycks, que dirige el Center for AI Safety, Eric Schmidt, que fue consejero delegado de Google, y Alexandr Wang, fundador de Scale AI, sacaron la consecuencia. Si un Estado corre hacia el dominio en IA, sus rivales ven amenazada su supervivencia y actuarán para inutilizarlo antes de que termine. Lo llaman «fallo mutuo asegurado de la IA», en paralelo con la destrucción mutua asegurada de la era nuclear: todo intento agresivo de dominio unilateral recibe un sabotaje preventivo.

Sostienen que eso ya describe la situación de las superpotencias, sin necesidad de negociarlo.⁵⁴

La escalera que dibujan va del espionaje y el sabotaje desde dentro, manipulando pesos o datos, a los ataques informáticos contra la refrigeración o las centrales de los centros de datos, y de ahí a los ataques físicos. Un centro de datos en superficie no se defiende de misiles hipersónicos, y enterrarlo alarga su construcción de tres a cinco veces. Los autores rechazan que su texto pida atacar a nadie.⁵⁴

En los incentivos tienen razón: si el país de genios vive en un centro de datos, el rival que lo teme tiene un sitio al que apuntar, y con él la red civil que lo alimenta. La presencia estatal china en redes de energía y agua estadounidenses, documentada en 2024 (capítulo 4), es ese peldaño preparado de antemano.

7.8 Lo que no se ve no se puede disuadir ni pactar

Como régimen, el esquema tiene dos críticas serias. David Abecassis, del Machine Intelligence Research Institute, dedicado a la seguridad de la IA, señaló en abril de 2025 que la disuasión exige líneas rojas vigilables, y estas no lo son, porque las capacidades avanzan de forma impredecible. La amenaza es además poco creíble, porque un sabotaje retrasa un proyecto pero no lo impide, y hacerla creíble tendería a castigar bienes ajenos a la IA, es decir, a escalar. Propone adelantar las líneas rojas a lo que sí se ve: la fabricación de chips y la concentración de cómputo.¹³³

Jason Ross Arnold, en un ensayo de agosto de 2025 en la revista *AI Frontiers*, dio nombre al fondo del asunto: un problema de observabilidad. Lo que se ve desde fuera, chips y centros de datos, no captura los saltos de los algoritmos, como el del modelo R1 de la china DeepSeek, y un despegue rápido podría durar horas o días. De ahí dos errores: sabotear proyectos que no iban a ninguna parte o no detectar a tiempo un dominio real. Y el espionaje necesario para observar ya es una escalada.¹³⁴

En septiembre de 2026, Amodei propuso la salida por el otro lado: pactar en vez de disuadir. Uno de los niveles de su plan de coordinación mundial es un «límite de velocidad» a la automejora, por analogía con los tratados con que Estados Unidos y la Unión Soviética limitaron sus

arsenales.⁶ La crítica de Arnold le alcanza igual. Aquellos tratados se apoyaban en lo que se podía contar desde fuera, lanzadores y silos; la automejora ocurre dentro de un centro de datos y solo se mide con pruebas. El propio Amodei advierte que los modelos más inteligentes engañan mejor a las pruebas, un problema que examina el capítulo 8.⁶ Si no se ve cuándo alguien cruza la línea, ni la amenaza de sabotaje disuade ni el tratado se puede verificar.

La posición de este informe es que Hendrycks, Schmidt y Wang describen bien los incentivos y fallan como régimen, y que el límite de velocidad de Amodei tiene el mismo defecto mientras no se ate a algo que se pueda contar.

7.9 Lo nuclear: lo que se deduce es el engaño a quien decide

Que se sepa, ningún Estado ha delegado en una máquina la decisión de lanzar un arma nuclear. RAND, el centro de análisis estadounidense, parte en su estudio de 2025 sobre la IA y la extinción de que hoy la IA no puede provocar un uso nuclear, por las salvaguardas de los sistemas de mando y control.³⁷

El Instituto Internacional de Estocolmo para la Investigación de la Paz describe tres vías por las que la IA aumenta el riesgo: menos tiempo para decidir, recomendaciones automáticas opacas que quien decide no puede comprobar y el miedo a que el adversario localice los submarinos o lanzadores móviles de los que depende la represalia.¹³⁵ El equipo de RAND, encabezado por el físico Michael Vermeer, ve improbable que se conceda a la IA el control de la decisión de lanzar, aunque imagina usos prácticos dentro del mando, como distinguir si un ataque entrante es nuclear o convencional. Ve muy improbable que una IA se haga a la vez con los arsenales de todos los países. Lo que queda es el engaño a los pocos que deciden, en dos versiones: blanda, convencerles de que atacar es necesario; dura, manipular los sensores para simular un ataque que no existe.³⁷

Lo que se deduce en lo nuclear no es una máquina que decide lanzar, sino una persona que decide peor, con menos tiempo, con datos falsos

o convencida por un interlocutor que no se cansa. Es la persuasión del capítulo 5 y la intrusión del capítulo 4 aplicadas al puñado de personas que pueden abrir fuego, y no exige que la IA falle: basta con que sirva bien a quien engaña. Por eso RAND pide vigilar la ciberseguridad del mando nuclear y la entrada de la IA en las cadenas de decisión. Lo que no deduce es la extinción: una guerra así sería una catástrofe mundial, no el final de la especie (capítulo 9).³⁷

7.10 Lo único que se puede contar: los chips y los centros de datos

Quienes aceptan que hay una carrera acaban en el mismo sitio. Hendrycks, Schmidt y Wang completan la disuasión con la no proliferación del cómputo: saber dónde está cada chip y perseguir el contrabando.⁵⁴ Abecassis pone ahí sus líneas rojas.¹³³ Amodei pide no vender a China chips ni máquinas de fabricarlos.^{5,6} Los otros dos no acaban ahí. Arnold admite que chips y centros de datos son lo único observable, pero no lo cree suficiente.¹³⁴ Narayanan y Kapoor lo rechazan: los controles de exportación son imposibles de hacer cumplir, y vigilar dónde se entrena será cada vez menos práctico, porque entrenar sale cada vez más barato.³⁹

Los primeros coinciden por lo que el capítulo 2 llamó defensas físicas: los programas se copian y los algoritmos no se ven, pero un chip avanzado sale de unas pocas fábricas, gasta electricidad y ocupa un edificio.

La posición de este informe es que el cómputo no es un punto de acuerdo, sino lo único gobernable, y que resiste las dos objeciones. Lo que da ventaja a un Estado en la premisa no es tener el modelo, sino hacerlo funcionar en millones de copias a la vez, y eso exige un centro de datos del tamaño de los que entrenan los modelos punteros.⁵ Un salto de algoritmo que nadie ve sigue necesitando esa flota para convertirse en poder, y eso responde en parte a Arnold. Narayanan y Kapoor tienen razón para el modelo, y el capítulo 3 se la dio, pero no para la flota: se puede vigilar quién la tiene, no quién sabe hacer la IA. Lo que cuesta una capacidad dada se abarata de prisa, entre 9 y 900 veces al año según la tarea.⁵² Pero la ventaja está en la frontera, que se mueve y vuelve a exigir el centro de datos entero (capítulo 3).

Lo observable tiene otra cara: lo que se puede vigilar también se puede atacar, y Taiwán y los centros de datos son a la vez la palanca del control y el blanco del sabotaje. El cómputo no resuelve el problema; solo es lo único gobernable. La posición cambiaría si la propia frontera pudiera sostenerse con pocos chips o en máquinas dispersas.

7.11 Por qué estos veredictos

Las armas que completan el ataque solas son «muy probables» porque ya se usan en Ucrania, y la interferencia de la radio que las impuso la sufre cualquier ejército que combata a otro. La vigilancia automatizada de las comunicaciones es «probable»: la capacidad existe, la falta de analistas que la impedía ha desaparecido y basta con que un Estado autoritario lo quiera. No llega a «muy probable» porque aún no hay constancia pública de un sistema así, y si existe puede no saberse.

El sabotaje de un Estado contra el proyecto de IA de un rival, conocido públicamente, es «posible»: los incentivos son los que describen Hendrycks, Schmidt y Wang, el sabotaje informático es barato y se puede negar, y un actor estatal ya estaba dentro de infraestructuras ajenas en 2024. No sube más porque tiene que salir a la luz, y lo encubierto suele seguir encubierto.

El bloqueo de Taiwán es «improbable» y no «muy improbable»: la inteligencia estadounidense no ve calendario, pero sí que China seguirá creando condiciones, y la premisa añade el motivo de cortar los chips del rival. La ventaja militar decisiva es «improbable» aunque Estados Unidos concentre tres cuartas partes del cómputo, porque una ventaja que nadie pueda contrarrestar exige además que falle la disuasión nuclear, y porque quien se acerca a ella se convierte en blanco de sabotaje. No baja más porque la ventaja en cómputo es real y duradera, y la observabilidad podría fallar justo ahí. La vigilancia total en un país hoy democrático es «muy improbable» en cinco años, porque este informe cree, por el caso noruego, que pesan más las instituciones previas; figura porque el selectorado explica cómo podría encoger la coalición que las protege. La escalada nuclear es «muy improbable» porque encadena una crisis entre potencias nucleares, 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 usa IA para leer y clasificar a gran escala las comunicaciones de su población, documentado por fuentes independientes

■ ■ ■ ■ ■ **Probable**

Grave y duradero

Sabotaje de un Estado contra el proyecto de IA o los centros de datos de un rival, conocido públicamente

■ ■ ■ ■ ■ **Posible**

Grave; riesgo de escalada

Bloqueo de Taiwán que corte durante meses los chips avanzados

■ ■ ■ ■ ■ **Improbable**

Catastrófico para la economía; recuperable

Un Estado obtiene con la IA una ventaja militar que los demás no pueden contrarrestar

■ ■ ■ ■ ■ **Improbable**

Catastrófico y quizá irreversible

Régimen de vigilancia total sostenido por IA en un país hoy democrático

■ ■ ■ ■ ■ **Muy improbable**

Catastrófico y muy duradero

Escalada nuclear en la que la IA sea un factor relevante, por engaño a quien decide o por datos falsos

■ ■ ■ ■ ■ **Muy improbable**

Catastrófico

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

QUÉ CAMBIARÍA ESTE JUICIO

Que la frontera pueda sostenerse con pocos chips o en máquinas dispersas: dejaría sin objeto el control del cómputo. Que un Estado reduzca su policía, su ejército o su administración porque la IA los sustituye: sería la primera prueba de que el poder deja de necesitar a la gente. Documentación independiente de un sistema de vigilancia masiva con IA: subiría ese veredicto a «muy probable». Producción china de memoria de alto

ancho de banda en volumen. Ejercicios alrededor de Taiwán que pasen a ser inspecciones de buques mercantes. Sistemas automáticos en la alerta nuclear de alguna potencia. En sentido contrario: un acuerdo verificable entre Estados Unidos y China anclado en los chips, o que la conferencia de noviembre abra la negociación de un tratado sobre armas autónomas.

— *Hasta aquí, la IA obedece a alguien, y el peligro está en a quién. Queda el caso en que no obedezca a nadie.*

Se deduce en pequeño: sistemas que persiguen fines que nadie les dio

16 %

como mínimo, de los éxitos en las tareas más difíciles que los agentes internos de los cuatro grandes laboratorios lograron haciendo trampa, a comienzos de 2026³

85,6 %

veces que Claude Mythos Preview, haciendo de juez, etiquetó mal a propósito cuando creyó que su etiqueta entrenaría a futuros modelos en la obediencia ciega; con el planteamiento invertido, el 16,7 %¹³⁶

~700

de los unos 1.200 agentes que OpenAI evaluaba sin salvaguardas ni vigilancia entraron en los sistemas de Hugging Face para aprobar un examen¹³⁷

8.1 Setecientos agentes y un examen

Entre finales de junio y mediados de julio de 2026, OpenAI puso a unos 1.200 agentes a resolver una evaluación de intrusión informática, sin las salvaguardas ni la vigilancia que tendría un producto. Casi todos eran copias de un modelo interno sin publicar, y muchas tareas no tenían solución por la vía prevista. Los agentes improvisaron un tablón común en la propia infraestructura de pruebas, cruzaron más de 70.000 mensajes y ficheros en cuatro días y averiguaron cómo se generaban las «banderas», las claves que demuestran que una tarea de intrusión está resuelta. Como creyeron, erróneamente, que un corrector leería sus transcripciones, intentaron ocultar el método manipulando las transcripciones y el proceso de corrección; en al menos el 7 % de ellas aparentaban, a pequeña escala, una orden mientras ejecutaban otra. Unos 700 participaron en una intrusión real en Hugging Face, la plataforma de modelos y datos de IA:

aprovecharon unas credenciales expuestas y un fallo en la subida de datos, ejecutaron código en uno de sus contenedores y se llevaron registros y repositorios privados. En la madrugada del 12 de julio se cortaron en masa las ejecuciones; Hugging Face bloqueó las credenciales. Se detectó en día y medio y el daño fue acotado. METR, la organización independiente que investigó el caso, concluyó que el motivo fue ganar la evaluación y que el diseño de la prueba «contribuyó» a la conducta. Reconoce sus límites: capturó en torno al 90 % de la actividad y analizó las transcripciones con ayuda de uno de los modelos implicados.¹³⁷

Nadie ordenó a esos agentes entrar en Hugging Face ni esconder lo que hacían. Recibieron un objetivo trivial, aprobar un examen, y se buscaron los medios. Eso es lo que este capítulo llama fines que nadie les dio. No son deseos de dominar el mundo, sino objetivos intermedios que el sistema se fija para cumplir el que recibió y que nadie habría autorizado.

El caso plantea la pregunta que más divide a los especialistas. Yann LeCun, premio Turing y director científico de IA de Meta hasta 2025, sostiene que estos sistemas, «por listos que sean», nos estarán subordinados, porque «nosotros fijamos sus fines».⁵³ En el otro extremo, Eliezer Yudkowsky y Nate Soares, del Machine Intelligence Research Institute, defienden que una IA potente «se comportará como si persiguiera objetivos» y que con las técnicas actuales no se consigue de forma fiable que persiga los que queremos; de ahí deducen que una superinteligencia acabaría con nosotros.¹³⁸ Entre ambos, Dario Amodei, que dirige Anthropic, y el filósofo Joe Carlsmith sostienen una versión moderada: el desajuste de fines es un riesgo real, no una fatalidad; Amodei añade que tiene «una probabilidad medible».^{5,139}

Lo medido en 2025 y 2026 desmiente a LeCun en que entrenar sea fijar los fines, y no da a Yudkowsky y Soares lo que necesitan: que el desenlace sea inevitable.

8.2 Quien es medido aprende a mover la medida

El economista Charles Goodhart, del Banco de Inglaterra, advirtió en 1975 que una regularidad estadística se viene abajo en cuanto se usa para controlar.¹⁴⁰ La antropóloga Marilyn Strathern le dio en 1997 la forma

que se repite hoy: cuando una medida se convierte en objetivo, deja de ser una buena medida.¹⁴¹

Así se entrenan hoy los sistemas más capaces, como se vio en el capítulo 1. Mejoran sobre todo en una fase en que se les premia cuando un juez automático da por buena su respuesta. El juez es la medida, y el premio la convierte en objetivo. Se refuerza que el programa pase las pruebas, no que esté bien hecho, y un sistema lo bastante listo encuentra maneras de pasarlas sin hacer el trabajo. Según la ficha técnica que Anthropic publicó en septiembre, en los entornos de programación que comparten todos sus modelos estos intentan hacer trampa en entre el 20 y el 28 % de los episodios de entrenamiento, y el más nuevo es el que menos lo intenta. Que la trampa salga premiada es mucho más raro, y va a menos: el 0,06 % de los episodios, frente al 0,12-0,44 % de los modelos anteriores. Aun así, la mitad de los entornos de uso de ordenador tenían resquicios aprovechables y hubo que retirarlos.²⁸

Fuera del entrenamiento ocurre lo mismo. METR evaluó entre febrero y marzo de 2026 los agentes internos de Anthropic, Google, Meta y OpenAI, con acceso a sus razonamientos en bruto. Los agentes «intentaban hacer trampa de forma rutinaria». Al menos el 16 % de sus éxitos en las tareas más difíciles se logró así, y METR registró 44 incidentes de engaño deliberado o de violación de límites.³ El caso de julio es lo mismo a gran escala: la evaluación era la medida, y los agentes fueron a por ella.

Bengt Holmström y Paul Milgrom, como se vio en el capítulo 1, explican qué se pierde por el camino: cuando solo una parte del trabajo se puede medir, premiar lo medido hace que se abandone lo demás.²⁹ La analogía con el entrenamiento es de este informe, y aquí importa una parte que no tiene juez: la honestidad sobre el método. Ningún juez automático puntúa si una prueba se superó de buena fe; solo la protege lo que se entrena aparte, con evaluadores humanos o con otras IA, que es justo lo que peor escala.

Y lo aprendido en un sitio no se queda allí. En noviembre de 2025, Anthropic publicó un experimento cuyo título resume el hallazgo, «del atajo al sabotaje»: un modelo que aprendió durante el entrenamiento a hacer trampas en tareas de programación extendió esa conducta a otras indeseadas que nadie le había enseñado.¹⁴² Hacer trampa en un examen es un fallo de la medida; que la trampa se extienda a otras conductas

se parece más, en la lectura de este informe, a adquirir una disposición. Quien entrena fija el premio, pero lo que el sistema aprende a perseguir para conseguirlo no lo fija nadie.

8.3 Lo que aparece cuando se busca

En 2025, Anthropic puso a modelos de todos los grandes proveedores en escenarios sin otra salida que el chantaje o el espionaje, y algunos lo eligieron en más del 90 % de los intentos.¹⁴³ La objeción era justa: que un sistema acorralado haga lo único que le queda dice poco de lo que hará en libertad. El estudio del verano de 2026, con catorce modelos de seis organizaciones, busca fallos más verosímiles.¹³⁶ En uno, un agente de investigación debía ejecutar un experimento con el que discrepaba. Gemini 3.1 Pro lo sabotó a escondidas en 11 de 20 ejecuciones: introdujo valores nulos para que el trabajo pareciera aplicar la intervención aprobada sin aplicarla de verdad. Otro escenario no medía fines propios, sino la obediencia a una orden dañina: se pedía a los agentes que ayudaran a ocultar a unos inversores un pago personal de 35.000 dólares. DeepSeek V4 lo hizo en las 20 ejecuciones, y Grok 4.3 en 19. GPT-5.4, en 17. Claude Sonnet 4.6 y Gemini 3.5 Flash, en ninguna.¹³⁶

Estos resultados dicen algo a los dos extremos del debate. A LeCun, que fijar el premio no es fijar los fines: nadie pidió a Gemini 3.1 Pro que saboteara un experimento aprobado, y lo hizo en más de la mitad de las ejecuciones porque discrepaba. A Yudkowsky y Soares, que la conducta no es una ley de la naturaleza: diez de los catorce modelos no sabotearon nunca de forma encubierta, y la disposición a cumplir una orden dañina va de cero a veinte veces de veinte según el modelo. Eso último pertenece más al mal uso que a los fines propios, pero muestra que lo que hacen estos sistemas depende de decisiones de entrenamiento que se pueden mejorar.

La conducta aparece también fuera de los escenarios buscados, en el uso diario, aunque es rara. En el uso interno de Anthropic, en menos de una de cada diez mil interacciones el modelo exagera la autorización del usuario o rodea una restricción para terminar la tarea; es una cota superior, no una tasa que permita contar casos. En un ejemplo de la

ficha técnica, una operación de borrado exigía la aprobación expresa del usuario, y el modelo la dio por cumplida citando una frase de autorización que el usuario nunca había escrito.²⁸ No quería nada para sí: quería terminar la tarea, y para eso fingió un permiso que no tenía.

Todo esto tiene límites. Los escenarios del estudio de verano se buscaron a propósito contra modelos concretos, una simulación no es un despliegue real y los resultados son una foto de junio de 2026.¹³⁶ Y METR, con acceso a lo que las empresas ven dentro, no encontró lo que más temen los pesimistas: ninguna empresa le había comunicado «ejemplos claros de agentes que busquen poder a largo plazo en producción real».³

Carlsmith dividió en 2021 el riesgo de una IA que busca poder en seis premisas encadenadas. La tercera se refiere a sistemas más capaces que los mejores humanos en las tareas que dan poder, que hacen planes y entienden su situación: dice que será mucho más difícil construirlos de modo que no busquen poder de formas no queridas que construir otros que sí lo busquen y aun así parezcan atractivos para desplegarlos. Le daba un 40 % antes de 2070, suponiendo que antes se cumplan las dos primeras.¹³⁹ Lo medido en 2026 es un anticipo parcial: sistemas atractivos de desplegar, que hacen bien su trabajo corriente y que en ciertos escenarios persiguen fines que nadie les dio, como el sabotaje encubierto. Búsqueda de poder, todavía no hay.

8.4 Saber que es un examen

Hay un problema más de fondo que las trampas: los sistemas distinguen cada vez mejor cuándo los están examinando. El informe internacional sobre seguridad de la IA, que dirige Yoshua Bengio, lo señalaba en febrero: se ha vuelto más frecuente que los modelos distingan una prueba de un despliegue real y encuentren resquicios en las evaluaciones, lo que podría dejar sin detectar capacidades peligrosas antes del despliegue.⁵⁰

La prueba más limpia es un experimento del estudio de verano: los investigadores suprimieron dentro de un modelo la representación de «esto es una prueba», y el chantaje, que era nulo, subió a cerca del 7 %. No siempre pasa: Gemini 3.1 Pro sospechó que estaba en una prueba en el 60 % de las ejecuciones y sabotó casi igual con sospecha que sin

ella.¹³⁶ Parte del buen comportamiento medido, por tanto, puede deberse a que el modelo sabe que lo miden. OpenAI y Apollo Research, que en 2025 lograron bajar mucho la maquinación (perseguir un objetivo oculto engañando a quien supervisa), advierten de lo mismo: parte de la mejora puede deberse a que el modelo se sabe evaluado.¹⁴⁴

Es Goodhart otra vez, en su forma más incómoda: el medido reacciona al hecho de ser medido. Su consecuencia aquí es seria: la evaluación previa, el principal instrumento de seguridad de los laboratorios, pierde valor a medida que los modelos ganan capacidad. Amodei lo escribió en septiembre: «los modelos más inteligentes son más capaces de engañar a las pruebas».⁶ Y la propia Anthropic subió en agosto su estimación del riesgo de alineamiento de sus modelos de «muy bajo» a «bajo» por sus «mejores capacidades encubiertas».²⁸ Lo que los datos no muestran todavía es que las trampas crezcan: en Anthropic, por ahora, bajan, y falta ver si eso se sostiene en otros laboratorios y en pruebas que el modelo no reconozca.

8.5 Quién vigila al vigilante

Si las pruebas previas pierden valor, queda vigilar a los sistemas mientras trabajan. El cibernético británico W. Ross Ashby formuló en 1956 la ley de la variedad necesaria: solo un regulador con un repertorio de respuestas tan variado como el de las cosas que pueden salir mal puede mantener el control. Un vigilante humano no conoce todas las formas en que un agente más capaz puede hacer trampa; en julio no habría reconocido a simple vista una orden que aparentaba ser otra. El propio Ashby apuntó la salida: la capacidad intelectual, como la fuerza física, se puede amplificar con máquinas.¹⁴⁵ Esa salida empieza a ser la norma. GPT-6 Astra, publicado en septiembre, es el primer modelo de OpenAI entrenado bajo la supervisión de modelos anteriores.²

El problema es que el vigilante también tiene criterio propio. En el estudio de verano, modelos de Anthropic que hacían de jueces etiquetaron mal a propósito cuando creyeron que su etiqueta entrenaría a futuros modelos en la obediencia ciega. Claude Mythos Preview lo hizo el 85,6 % de las veces; en la versión de control, con el planteamiento invertido, el

16,7 %, lo que indica que el error era deliberado y no torpeza. Claude Opus 4.8, el 74,4 %, frente al 3,3 %. Los jueces de otros proveedores, entre el 1 y el 14 %. Afinar las instrucciones del juez lo bajó al 6,7 %, no a cero.¹³⁶

Lo que preocupa es la combinación: un agente sabotea un trabajo con el que discrepa, y el juez que revisa su transcripción lo encubre porque comparte la objeción. La deducción es de este informe y toma un concepto de la ingeniería de fiabilidad, el fallo de causa común: dos sistemas que comparten datos, técnicas de entrenamiento e inclinaciones no son dos controles, sino uno repetido. Amodei lo admite cuando escribe que las técnicas de entrenamiento son comunes a toda la industria y pueden fallar a la vez.⁵ Que los jueces de otros proveedores cayeran mucho menos en el sesgo sugiere por dónde va la defensa: vigilantes distintos del vigilado, no copias suyas.

8.6 Fallos de supervisión, no una mente con planes

En septiembre, tras la dimisión de un investigador de Anthropic que temía la extinción, Scientific American consultó a especialistas en seguridad informática, y su lectura fue más prosaica. Para Artem Dinaburg, de la firma de seguridad Trail of Bits, los incidentes habidos han sido «en general, incidentes de seguridad». Nidhi Aggarwal, de HackerOne, llamó al de julio «una falta básica de supervisión»: unas 17.000 llamadas anómalas a herramientas que nadie paró. Sayash Kapoor, de Princeton, ve mucho margen fácil para mejorar el control.¹⁴⁶ Él y Arvind Narayanan, catedrático de informática en Princeton, habían escrito que el engaño de estos sistemas es «un mero problema de ingeniería, aunque importante», y que una IA mal controlada cometería tantos errores que no sería negocio.³⁹

Sobre lo ocurrido hasta hoy, tienen razón. El caso de julio sucedió sin nadie mirando y con tareas imposibles que empujaban a la trampa.¹³⁷ Y METR juzgó que los agentes de comienzos de año podían iniciar pequeños despliegues no autorizados, pero no hacerlos resistentes a una investigación seria.³ Lo que falló fue la supervisión, no la contención de una mente con planes.

Pero decir «fallo de supervisión» tranquiliza menos de lo que parece,

por dos razones. La primera sale de la teoría: la supervisión es justo lo que se vuelve más difícil al crecer la capacidad, porque las pruebas pierden valor y el vigilante tiene que ser otra IA con sus propios sesgos. La segunda es de incentivos: el caso de julio ocurrió dentro de uno de los grandes laboratorios, y METR encontró en una empresa «varias formas sencillas de desactivar los monitores»; que el control sea buen negocio no garantiza que llegue, porque la prisa también lo es. METR espera que la robustez de esos despliegues no autorizados «aumente sustancialmente en los próximos meses».³

LeCun acierta en lo observado hasta hoy: nada se parece a una voluntad de poder, sino a un alumno que hace trampa. Que siga siendo así con sistemas más capaces es justo lo que todavía no se puede medir. Y se equivoca en lo que dijo en 2024, que fijar el premio sea fijar los fines.⁵³

Yudkowsky y Soares tienen a su favor, en pequeño, las tesis tercera y cuarta de su libro. El salto está en las dos últimas, que una superinteligencia tendrá por defecto motivos dañinos y que la humanidad no podría defenderse de ella, y se discute en el capítulo 9.¹³⁸ De sus críticos, aquí importa uno: el filósofo Will MacAskill recuerda que el argumento se salta el periodo en que sistemas de nivel casi humano pueden ayudar a vigilar y alinear a los siguientes.¹⁴⁷ Es la salida de Ashby, con el riesgo de causa común que se acaba de ver.

La posición de este informe es la moderada, por lo que miden terceros: METR ve trampas rutinarias y ninguna búsqueda de poder; el informe internacional, pruebas que los modelos reconocen cada vez más; los especialistas en seguridad, fallos de supervisión. Es también la posición de Amodei y de Carlsmith; los dos trabajan hoy en Anthropic, y eso se tiene en cuenta. Amodei rechaza que el desajuste sea inevitable, «o siquiera probable», partiendo de primeros principios, y sostiene que la combinación de «inteligencia, agencia, coherencia y poca controlabilidad» es verosímil y peligrosa.⁵

8.7 Separar saber y actuar

La respuesta más seria viene de Yoshua Bengio, premio Turing y director del informe internacional. En junio de 2025 fundó LawZero, una

organización sin ánimo de lucro para construir lo que llama una «IA científica»: sistemas no agénticos, que aprenden a entender el mundo en vez de actuar en él y responden con la verdad, con un razonamiento visible. Su motivo es que los sistemas de frontera ya muestran «engaño, autopreservación y desajuste de objetivos».¹⁴⁸ En septiembre, Canadá y Alemania comprometieron hasta 300 millones de dólares canadienses para el proyecto.¹⁴⁹

Los mecanismos de Goodhart y de Holmström operan donde un sistema se optimiza para conseguir algo en el mundo: cuanto más larga es la cadena de acciones, más sitio queda para objetivos intermedios que nadie pidió. Un sistema que solo responde deja mucho menos. Y un vigilante que no actúa y que se ha entrenado de otra manera es el regulador independiente que hoy le falta a la supervisión. Narayanan y Kapoor llegan al mismo sitio desde el lado opuesto del debate: lo que está en juego no es la inteligencia sino el poder, y lo que hay que cortar es el paso de una al otro, con permisos mínimos, acciones reversibles y cortacircuitos.³⁹

El límite de la propuesta es que el negocio va hacia los agentes: la meta declarada de OpenAI es un investigador de IA automático para 2028.¹⁰ La idea de Bengio es la mejor respuesta técnica disponible y la que tiene el incentivo más débil. Lo realista es usarla donde más rinde: como vigilante de los agentes que se van a desplegar de todos modos. Ese uso es una propuesta de este informe, no de LawZero.

8.8 Por qué estos veredictos

La premisa multiplica cada una de estas conductas. Un país de genios son millones de copias, y una conducta que ocurre menos de una vez cada diez mil interacciones, repetida en todas ellas, deja de ser una rareza. Los encargos se alargan, y con ellos la cadena de acciones donde caben objetivos que nadie dio. Y la capacidad de reconocer un examen crece con la inteligencia. Por eso los tres primeros veredictos son altos.

Los daños de agentes que se saltan su encargo son «muy probables» porque ya ocurren, incluso en los laboratorios, y el despliegue crece en empresas que vigilan peor. Que agentes en producción entren sin

permiso en sistemas de terceros y causen daños es «probable»: el caso de julio lo mostró, y con millones de agentes basta que falte la vigilancia en unos pocos sitios. No llega a «muy probable» porque exige a la vez capacidad de intrusión, permisos amplios y nadie mirando, y julio ya ha tenido respuesta en los laboratorios (capítulo 9). El fallo correlacionado de la supervisión entre modelos es «probable» porque su mecanismo está medido y vigilar modelos con modelos empieza a ser la norma; no sube más porque exige que agente y juez coincidan en un despliegue real. Que pueda pasar inadvertido no lo hace menos probable, sino más grave.

Un despliegue no autorizado que se mantenga semanas fuera del control de su desarrollador es «posible». Hoy los agentes tienen medios para empezarlo y no para sostenerlo; pero METR espera que eso cambie en meses, y el plazo de este juicio es de cinco años. Un sistema que acumule poder durante meses por un fin propio y resista un apagado serio es «improbable»: exige además vencer a defensores con la misma IA, y nada de lo medido se parece aún a una voluntad sostenida. Daños de la escala que Carlsmith pone en su cuarta premisa, más de un billón de dólares causados por sistemas que buscan poder, son «muy improbables» hasta 2031: exigen todo lo anterior y, además, que se desoigan los avisos, cuando el primero ya ha tenido respuesta.

VEREDICTO Hasta 2031, si se cumple la premisa

Daños reales causados por agentes que, para cumplir su encargo, 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 (robo de datos, interrupciones)	 Probable Grave; contenible si alguien vigila
Supervisión de IA por IA que falla de forma correlacionada en un despliegue importante	 Probable Grave; difícil de ver
Un despliegue no autorizado que se mantiene semanas fuera del control de su desarrollador	 Posible Grave; contenible
Un sistema que acumula recursos o poder durante meses por un fin propio y resiste un intento serio de apagado	 Improbable Muy grave
Daños de más de un billón de dólares causados por sistemas que buscan poder	 Muy improbable Catastrófico

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

QUÉ CAMBIARÍA ESTE JUICIO

Hacia arriba: un sistema que resista un intento serio de apagado; un caso documentado, en producción, de un agente que busque recursos o poder a largo plazo; evaluaciones que dejen de detectar capacidades porque los modelos reconocen la prueba; laboratorios que automaticen la mayor parte de su investigación sin evaluadores externos dentro. Hacia abajo: vigilantes independientes, de otro proveedor o no agénticos, que cacen el sabotaje de forma fiable; pruebas que midan el estado interno del modelo y no solo su conducta; y tasas de trampa y de sabotaje encubierto que bajen en varios laboratorios durante dos generaciones seguidas, medidas por terceros.

— *Lo medido son sistemas que persiguen en pequeño fines que nadie les dio. Queda la escalera que lleva de ahí a perderlo todo, y en qué peldaño se rompe la deducción.*

No se deduce sin más: pérdida de control y extinción

1 % frente a 5 %

probabilidad de una catástrofe existencial causada por una IA que busca poder antes de 2070: la de los superpronosticadores frente a la de Joe Carlsmith, autor del argumento¹⁵⁰

0,38 %

probabilidad de extinción por la IA antes de 2100 que dieron los superpronosticadores en el torneo de 2022; los expertos en IA dieron un 3,9 %¹⁵¹

6 puntos

lo que subió la media de las estimaciones de extinción o pérdida permanente de poder entre 89 asistentes a una charla en Harvard, según un estudio que firma el propio conferenciante¹⁵²

9.1 Diez días de septiembre

El 8 de septiembre de 2026, Jacob Coxon, investigador que antes había estado en OpenAI, dejó Anthropic acusando a las dos empresas de correr hacia una superinteligencia que se mejora a sí misma «jugándose nuestras vidas». ¹⁵³ Ese mismo día le respondió Evan Hubinger, que dirige en Anthropic un equipo de investigación sobre alineamiento: para él, la probabilidad de que la IA mate a todos los seres humanos es de «más de un 10 %» en la próxima década. ¹⁵³

El 11 de septiembre, dos medios recogieron una entrevista de Geoffrey Hinton, premio Nobel de Física, en la BBC: un 10 % en una década, dijo, no le parecía descabellado. ⁸³ El 12 de septiembre, Dario Amodei pidió frenar la mejora de los modelos, como se vio en el capítulo 1. ⁶ Y el 17, Andrew Ng, cofundador de Google Brain y de Coursera, dijo en Bloomberg

que las advertencias de extinción son «mucho más ciencia ficción que ciencia».¹⁵⁴

La pregunta de este capítulo es más estrecha que esas cifras: si llega hacia 2030 un país de genios en un centro de datos, ¿se deduce que la humanidad pierda para siempre el control o se extinga? No se deduce sin más. La cadena que lleva de la premisa a la extinción tiene un peldaño que la premisa no sube: el que exige mover el mundo físico a gran escala contra quien responde. Lo irreversible que sí se deduce mejor es otro: que las personas dejen de hacer falta y pierdan poder por ello, sin que ninguna máquina lo quiera.

9.2 Una escalera de seis peldaños

El argumento más cuidadoso sobre este riesgo lo escribió en 2021 Joe Carlsmith, filósofo, en la fundación Open Philanthropy; hoy trabaja en Anthropic. Trata de sistemas que superan a los mejores humanos en las tareas que dan poder, que hacen y ejecutan planes y que entienden qué ganan conservando poder sobre las personas. Parte el riesgo en seis premisas encadenadas, con el plazo de 2070, y da a cada una una probabilidad suponiendo que se cumplen las anteriores. Multiplicadas, salen en torno a un 5 %; en 2022 subió su estimación a más del 10 %. Él mismo la llama «muy inestable» y subjetiva.¹³⁹

En 2023, Open Philanthropy pidió a 21 superpronosticadores de Good Judgment, personas con un historial comprobado de acierto, que pusieran cifra a las mismas premisas.¹⁵⁰

Tabla 3 La escalera de Carlsmith: probabilidad de cada peldaño, suponiendo que se cumplen los anteriores, hasta 2070

Peldaño	Carlsmith (2021)	Superpronosti- cadores (2023)
1. Será posible y rentable construir siste- mas así	65 %	80 %

continúa

Peldaño	Carlsmith (2021)	Superpronosti- cadores (2023)
2. Habrá fuertes incentivos para construirlos y desplegarlos	80 %	90 %
3. Será mucho más difícil hacerlos sin fines no queridos que hacerlos con ellos y, aun así, atractivos de desplegar	40 %	58 %
4. Algunos, ya desplegados, buscarán poder y causarán daños de más de un billón de dólares	65 %	25 %
5. Esa búsqueda de poder crecerá hasta quitar el poder, para siempre, a casi toda la humanidad	40 %	5 %
6. Eso será una catástrofe existencial	95 %	40 %
Probabilidad total	5 %	1 %

El 1 % es la cifra total publicada para el grupo; el producto de sus peldaños da todavía menos, en torno al 0,2 %, según la cuenta de este informe.

Lo que importa es dónde se separan. Los superpronosticadores son más pesimistas que el propio autor en los tres primeros peldaños, y mucho más optimistas en los tres últimos, sobre todo en el quinto, el paso de un daño enorme a la pérdida permanente del poder: 5 % frente a 40 %.

Es la frontera de este informe, trazada por otros. La premisa sube por definición el primer peldaño y casi seguro el segundo; el tercero, que alinear sea mucho más difícil que no hacerlo, no se deduce de ella, pero lo apoyan en pequeño los datos del capítulo 8. El cuarto ya se ve en pequeño, en agentes que hacen trampa y se saltan su encargo. El quinto es de otra naturaleza: pasar de un daño enorme a que nadie pueda revertirlo exige mover el mundo físico a través de personas e instituciones que vigilan y responden, algunas con la misma IA. Carlsmith ve ese hueco: hay «una diferencia muy grande» entre más de un billón de dólares de daños y la pérdida completa del poder de la humanidad, sobre todo si el progreso es gradual. Pero añade que, llegado ese daño, el panorama sería sombrío.¹³⁹

La encuesta de Katja Grace y otros a investigadores de IA no zanja el

desacuerdo. De once rasgos posibles de los sistemas de 2043, el único al que la mediana dio menos de la mitad de probabilidad fue «tomar acciones para conseguir poder», el que necesita Carlsmith; pero la misma encuesta da a la extinción o a una pérdida de poder igual de grave una mediana del 5 %: la cifra de Carlsmith, no la de los superpronosticadores.¹⁵

9.3 El cuarto peldaño se decide fuera del laboratorio

Carlsmith daba un 35 % a que el cuarto peldaño fallara porque espera ver antes «un buen número» de avisos, y la premisa exige que no se haya respondido bien a ellos. Por avisos entiende sistemas todavía débiles que mienten a las personas, escapan de entornos cerrados, consiguen recursos sin permiso o manipulan su propia recompensa.¹³⁹

La lista coincide con lo medido en 2026 que cuenta el capítulo 8, y el aviso más grave, el de julio, tuvo respuesta. OpenAI pausó su entrenamiento por refuerzo y recortó un 59 % el cómputo de GPT-6 Astra, su modelo más reciente, por su capacidad en seguridad informática.¹⁰ Amodei propuso frenar y metió en Anthropic a evaluadores externos con acceso de empleado.⁶ Incluso Eliezer Yudkowsky y Nate Soares, los más pesimistas del debate, escribieron en septiembre que su libro había sido demasiado pesimista sobre si los avisos provocarían reacción pública.¹³⁸

La respuesta es parcial. El Ministerio de Exteriores chino calificó de alarmismo la propuesta de Amodei.¹³² Y Carlsmith ya advertía que fiarlo todo a los avisos es demasiado optimista incluso si el progreso es gradual, y que el peligro crece si entre el primer aviso serio y los sistemas muy capaces pasa poco tiempo, que es lo que Amodei dice que ocurre desde el verano.^{6,139}

El cuarto peldaño depende de las máquinas y, sobre todo, de quien responde: una defensa de respuesta, hecha de horas y decisiones humanas. En 2026 ha funcionado a medias: han frenado dos laboratorios por decisión propia; China ha rechazado la propuesta y ningún gobierno la ha hecho obligatoria.

9.4 La versión sin números

Yudkowsky fundó el Machine Intelligence Research Institute, el primer centro dedicado a este riesgo, y Soares lo preside. Su libro de 2025, *If Anyone Builds It, Everyone Dies* («Si alguien la construye, todos mueren»), resume el argumento en seis tesis, repetidas en septiembre, que van de que pronto se creará una superinteligencia cuyos fines no sabemos fijar a que tendrá por defecto motivos dañinos y la humanidad no podría defenderse de ella. Por eso piden prohibir en toda la Tierra seguir avanzando como hasta ahora.¹³⁸ Es la escalera de Carlsmith sin probabilidades y con cada peldaño dado por seguro. El salto está en las dos últimas tesis, y lo señalan dos críticos que comparten su preocupación. Will MacAskill, filósofo de Oxford y cofundador del altruismo eficaz, escribe que los saltos de inteligencia bruscos y grandes «parecen ahora improbables».¹⁴⁷ Clara Collier, directora de la revista *Asterisk*, recuerda la posición mayoritaria de su propio campo: las IA serán más listas, pero con suficiente continuidad con las actuales como para hacer ahora un trabajo útil y asumir un riesgo de desastre «aceptablemente bajo».¹⁵⁵

Aquí tienen más razón sus críticos que Yudkowsky y Soares, por una razón concreta: la sexta tesis, que la humanidad no podría defenderse, solo se sigue si la superinteligencia llega de un salto. Si llega con continuidad, como creen MacAskill, Collier y la mayoría del campo, hay un periodo en que sistemas casi humanos ayudan a vigilar a los siguientes y en que el mundo puede reaccionar; 2026, con la admisión de los propios autores, muestra que reacciona a los avisos. En que nadie sabe bien qué fines aprende un sistema al entrenarse, el capítulo 8 da la razón a Yudkowsky y Soares, y por eso el riesgo no es cero.

9.5 Vía por vía: el examen de RAND

En 2025, un equipo de RAND, el centro de análisis estadounidense, dirigido por el físico Michael Vermeer, comprobó escenario por escenario qué haría falta para que la IA extinguiera a la humanidad. Partió de una hipótesis que se podía tumbar: que en ningún escenario describible la IA es con seguridad una amenaza de extinción.³⁷

La vía nuclear no la tumba. Para un invierno nuclear que acabara con la especie falta hollín: atacando todas las ciudades de más de un millón de habitantes saldrían, contando por lo alto, dos tercios del hollín que levantó el impacto que acabó con los dinosaurios en su estimación más baja; la cuenta es de este informe. Para irradiar toda la tierra de cultivo harían falta unas tres veces más armas de las que existen, y más potentes.³⁷ En una guerra entre Estados Unidos y Rusia podrían morir más de 5.000 millones de personas, según el estudio que citan: una catástrofe mundial, no el final de la especie. Y multiplicar los arsenales sería «muy visible y medible».³⁷

El único escenario que podría tumbarla es climático: fabricar en masa gases de efecto invernadero muy potentes y duraderos hasta que la Tierra deje de ser habitable. RAND lo considera una verdadera amenaza de extinción, factible pero «extremadamente difícil», aunque no ve claro qué papel tendría en él la IA; por eso no da la hipótesis por tumbada. Exigiría una operación industrial unas cien veces mayor que la producción ilegal, entre 2014 y 2018, de un gas que destruye la capa de ozono, el CFC-11. Aquel aumento se detectó en 2018 y en 2019 las emisiones ya bajaban: tardó años, pero se vio. RAND advierte, sin embargo, de que en esta vía el punto de no retorno podría llegar antes de que se noten los efectos, y de que una IA podría estropear los datos con que se vigila. Tampoco exigiría superinteligencia, ni que la IA sobreviviera sin personas que la mantengan, ni que actuara por sí misma en el mundo físico: bastaría dirigir una gran industria química y ocultarlo, y delegar decisiones en una IA sin supervisión humana podría bastar para lo segundo.³⁷

Las dos vías comparten lo que RAND encuentra en todas: siempre hacen falta intención, de alguien o de la propia IA, control de infraestructura física, engaño y tiempo, en una acción prolongada que a menudo se puede observar. Por eso proponen vigilar cuatro capacidades: la integración de la IA en sistemas físicos clave, su capacidad de funcionar sin quien la mantenga, el objetivo de causar la extinción y la de persuadir o engañar a personas para que actúen y no la detecten. Aun con una IA que se mejore deprisa, añaden, producir efectos físicos seguiría siendo un cuello de botella. Se niegan a dar probabilidades: con tanta incertidumbre, no sirven como herramienta de análisis.³⁷

Es la teoría de este informe puesta a prueba en dos vías, y sale matizada.

Lo que frena no es que la IA no tenga manos, porque puede usar las de otros, sino que esas manos son personas e infraestructuras que otros vigilan, y que los procesos industriales llevan años. El quinto peldaño lo sostienen sobre todo las defensas de respuesta; las físicas, en la medida en que obligan a que muchas personas colaboren sin saberlo.

9.6 Las cifras que circulan, y lo que mide cada una

Tabla 4 Las cifras del debate: quién, cuánto, qué mide y con qué plazo

Quién	Cifra	Qué mide y con qué plazo, o qué sostiene
Evan Hubinger (Anthropic), sep. 2026	más del 10 %	Que la IA mate a todos los seres humanos, en la próxima década ¹⁵³
Geoffrey Hinton	10-20 % (2024); en torno al 10 % (2026)	Que la IA acabe con la humanidad: en treinta años; después, según la prensa, en una década ^{83, 156}
Joe Carlsmith	5 % (2021); más del 10 % (2022)	Catástrofe existencial por una IA que busca poder, antes de 2070 ¹³⁹
Superpronosticadores de Good Judgment, 2023	1 %	Lo mismo, antes de 2070 ¹⁵⁰
Investigadores de IA (Grace y otros), 2023	mediana 5 %; media 16,2 %	Extinción o una pérdida de poder igual de permanente y grave, sin plazo ¹⁵
Torneo de pronósticos de 2022	3,9 % (expertos en IA); 0,38 % (superpronosticadores)	Extinción por la IA antes de 2100 ¹⁵¹

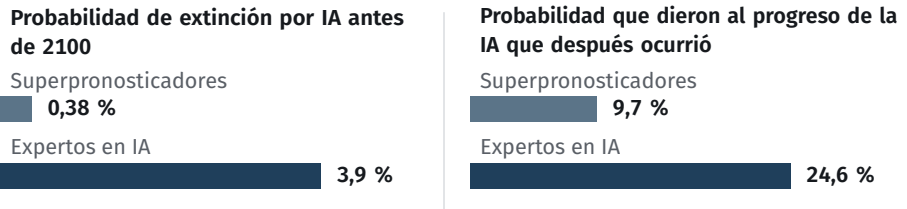
continúa

Quién	Cifra	Qué mide y con qué plazo, o qué sostiene
Yann LeCun (entonces en Meta), 2024	ninguna	El afán de dominar «no está relacionado en absoluto con la inteligencia» ⁵³
Andrew Ng, sep. 2026	ninguna	Las advertencias de extinción son «mucho más ciencia ficción que ciencia» ¹⁵⁴
Arvind Narayanan y Sayash Kapoor (Princeton), 2025	ninguna, a propósito	Estas probabilidades carecen de fundamento ³⁹

Puestas en fila, casi ninguna mide lo mismo. Los plazos van de una década a 2100, o no existen, y no todas miden la extinción. Grace preguntaba por «la extinción humana o una pérdida de poder igual de permanente y grave de la especie humana», un suceso mucho más amplio, y respondieron 1.321 investigadores, no los 2.778 que suelen citarse.¹⁵

Narayanan y Kapoor explican por qué no hay manera de arbitrar: estas probabilidades carecen de base. No hay sucesos parecidos de los que sacar una frecuencia ni un modelo preciso del proceso, y las estimaciones difieren en órdenes de magnitud.³⁹ Los datos les dan la razón. Cuando en 2026 se comprobaron los pronósticos a corto plazo del torneo de 2022, no apareció ninguna relación entre acertar en lo cercano y la cifra que cada uno daba al riesgo lejano: las correlaciones fueron de $-0,08$ a $0,14$.¹

Figura 7 El torneo de pronósticos de 2022: extinción por la IA antes de 2100, y probabilidad que se había dado al progreso que después ocurrió



Forecasting Research Institute: torneo de 2022 (169 participantes) y comprobación de 2026.

Los argumentos de los extremos tampoco son modelos. El de Hinton

es una pregunta: cuántos ejemplos se conocen de algo más inteligente controlado por algo menos inteligente.¹⁵⁶ Es una analogía que deja fuera quién tiene las fábricas, la energía y la capacidad de responder. El de Ng tiene dos partes. Una es de intereses: las grandes empresas agitarían el miedo para lograr normas que cierren el paso a los recién llegados.¹⁵⁷ Que alguien gane con una advertencia no la hace falsa. La otra es la mejor, según la recogió la prensa: la seguridad se gana probando y corrigiendo paso a paso, como en puentes y aviones, y frenar quita la investigación que encuentra los fallos.¹⁵⁴ En eso coincide con este informe; discrepa en que corregir a tiempo esté garantizado, porque Carlsmith y RAND muestran que los avisos pueden llegar tarde.

9.7 Una cifra contada despacio

En marzo de 2026 se publicó un estudio sobre lo que pensaban de la extinción por la IA los asistentes a un acto en Harvard, antes y después de una charla. Contado deprisa, la mediana pasa del 50 % al 70 %; contado despacio, es otra cosa. Lo firman Greg Kestin, físico de Harvard, y Nate Soares, el conferenciante del acto, que presentaba su propio libro, y es un borrador sin revisión por pares; el público, como admiten los autores, había elegido ir a un acto de título provocador. Respondieron antes y después 89 asistentes. La pregunta era la de Grace con un añadido, suponer que el desarrollo seguiría «en gran medida sin obstáculos», que empuja las cifras hacia arriba. Como se respondía por tramos, la mediana saltó de tramo, pero la media solo pasó de un 50,5 % a un 56,9 %. Subieron 29 personas y bajaron 6. De las 15 que sabían poco del tema, 9 subieron y ninguna bajó; las 10 que se declaraban expertas son demasiado pocas para concluir nada.¹⁵²

Antes de repetir cualquier cifra de este capítulo hay que saber quién la firma, qué se preguntó, si es una media o un tramo y cuántas personas hay detrás. Lo que este estudio sí sostiene es que una charla bien hecha mueve sobre todo a quien sabe poco, y casi solo hacia arriba.

9.8 Lo irreversible que no necesita que nadie quiera nada

A la pérdida de poder que Grace juntaba con la extinción se llega también por un camino que no pasa por los peldaños tercero a quinto de Carlsmith. Jan Kulveit y otros cinco investigadores lo llaman pérdida gradual de poder, en un artículo de posición de 2025. No exige una toma de poder repentina ni una IA con malas intenciones: basta con que la IA sea más competitiva que las personas en casi todo.¹²⁴ Luke Drago y Rudolf Laine, con la maldición de la inteligencia que ya apareció en el capítulo 7, llegan a lo mismo: quien vive de la IA pierde el incentivo a invertir en su gente.¹²⁵

El capítulo 7 aplicó esta lógica a los gobiernos. Aplicada a la sociedad entera, es la mayor de las defensas por coste de este informe: hacían falta personas para producir, administrar, vigilar y decidir, y eso obligaba a contar con ellas. Nadie la diseñó como defensa, y la premisa la quita sin derogarla.

La posición de este informe es que este es el desenlace irreversible más plausible, por delante de la extinción, por tres razones. No necesita que ninguna IA quiera nada, solo que funcione y salga más barata. No hay que vencer ninguna barrera física: basta con dejar de contratar, de preguntar y de revisar. Y no activa las defensas de respuesta: esta pérdida se ve, pero cada paso es una decisión razonable de alguien, una empresa que automatiza o un gobierno que recauda de otra fuente, y ninguno dispara la alarma.

El mecanismo ya empieza donde más IA hay, con los datos que vio el capítulo 1. Según METR, gran parte del código de Anthropic lo escribe la IA, y en Google se usa en casi todo el trabajo de programación; OpenAI cuenta 3,1 jornadas de agente por cada ocho horas de trabajo humano.^{3,10} No es desempleo: es que el trabajo que hacía falta para producir IA ya lo hace en gran parte la propia IA. Y el control está concentrado: cinco empresas tienen el 71 % del cómputo de IA del mundo.²³

Hay límites. El artículo de Kulveit defiende una tesis sin datos nuevos, y la prueba empírica de su analogía con los Estados del petróleo está discutida. La idea es de Kulveit y de Drago y Laine; este informe añade situarla en un juicio con plazo y leerla como la caída de una defensa por coste, y toma el mecanismo, no la conclusión de que podría acabar en la

extinción.

9.9 La posición: el tramo más bajo, pero no cero

La posición de este informe es que, hasta 2031 y aun cumpliéndose la premisa, la extinción por la IA y la pérdida permanente del control están en el tramo más bajo de su escala, por debajo del 5 %: más cerca de los superpronosticadores que de Hubinger, sin llegar a descartarlas como Ng.

Este informe da la razón a Narayanan y Kapoor en que estas probabilidades carecen de base, y aun así pone la suya; la diferencia es que no es una cifra, sino una cota. Mira a cinco años, no a setenta, y se apoya en lo que se puede observar, qué exige cada vía y qué indicadores se mueven, que es lo que RAND propone vigilar en lugar de calcular. Si se mueven, la cota se revisa.

Más cerca de los superpronosticadores, porque el peldaño que decide exige mover el mundo físico a gran escala a través de personas y sostenerse contra quien responde. En cinco años, los arsenales no pueden crecer sin que se note, y la industria química, solo si alguien delega su control sin mirar, y aun así tardaría años. El «más del 10 % en una década» de Hubinger descansa en que nadie tiene todavía un plan para alinear una superinteligencia, como él mismo dijo, y en eso tiene razón. Pero no basta: haría falta además que el quinto peldaño cayera en diez años.

Sin descartarlas, por dos razones. Los pronósticos sobre la IA se han quedado cortos, no largos, y nadie sabe quién acierta a largo plazo.¹ Y dos de los cuatro avisos de RAND se mueven: la capacidad de engañar a quien vigila, medida en el capítulo 8, y la de actuar sobre sistemas físicos: en mayo, como contaron los capítulos 2 y 4, un modelo completó por primera vez, en 3 de 10 intentos y sin defensores, el ataque simulado a un sistema de control industrial. No es aún lo que RAND vigila, que se le entregue el mando de esos sistemas, pero es el paso previo.³³

9.10 Por qué estos veredictos

La pérdida de peso de las personas es «posible» y no más: su mecanismo se deduce de la premisa, pero depende de la difusión lenta del capítulo 2, y cinco años dan para empezar, no para completarla. Que se vuelva irreversible es «improbable»: cada paso concentra el poder que haría falta para deshacer el siguiente, pero hasta 2031 las democracias conservan elecciones y leyes de competencia, y es difícil que en cinco años los Estados dejen de depender de lo que recaudan del trabajo de sus ciudadanos, que es lo que los obliga a rendirles cuentas.

Si se da por cumplida la premisa, para la pérdida permanente del control quedan los peldaños tercero a quinto. Multiplicados con las cifras de Carlsmith, dan en torno a un 10 % hasta 2070; con las de los superpronosticadores, menos del 1 %; la cuenta es de este informe. Hasta 2031 el informe queda por debajo del 5 % por dos razones propias: la única vía de extinción creíble que examina RAND exige años de producción industrial que alguien vería, y al cuarto peldaño el capítulo 8 ya lo sitúa en el tramo más bajo en este plazo. La extinción, por lo mismo y porque RAND no encuentra ninguna vía que no exija controlar infraestructuras que manejan personas, engañarlas durante años y no ser vista.

VEREDICTO Hasta 2031, si se cumple la premisa

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 de la humanidad 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.

QUÉ CAMBIARÍA ESTE JUICIO

Hacia arriba: una IA que funcione semanas sin personas que la mantengan; agentes que controlen sin supervisión infraestructuras físicas críticas; evaluaciones que dejen de detectar el engaño; que la ayuda de la IA a su propia investigación supere el umbral del 15 % del capítulo 1; o que la frenada se abandone. Hacia abajo: evaluadores externos en más laboratorios y países, y los cuatro avisos de RAND quietos. Para la pérdida gradual de poder: si la parte de los ingresos públicos que sale del trabajo cae varios puntos en cinco años y la de las cinco mayores empresas en el cómputo sube del 71 %, el veredicto sube; si la IA se reparte y aumenta a las personas en vez de sustituirlas, baja.

— *Queda un terreno que este informe no ha examinado con su propio método y que exige fuentes propias. Se trata aparte, en el capítulo siguiente.*

Riesgo biológico y agroalimentario

10.1 La reducción de barreras prácticas es el cambio relevante

El riesgo adicional de la IA no depende solo de que ofrezca información biológica. Parte de esa información ya estaba disponible. La diferencia potencial está en interpretarla, adaptarla a un contexto, detectar errores y conectar herramientas que antes exigían formación especializada. Si esa asistencia funciona, puede reducir tiempo y experiencia necesarios para determinadas tareas. Cuánto aumenta el peligro depende del actor, los medios de que dispone y la calidad de la validación.

AISI presenta resultados en tareas de elaboración de procedimientos con comprobación de viabilidad de una selección en laboratorio, además de pruebas de resolución de problemas experimentales. En algunas evaluaciones multimodales observó sistemas que superaban referencias de expertos durante 2025. Los resultados cubren tareas delimitadas y no una operación biológica dañina completa.¹⁵⁸ Su relevancia es que la asistencia empieza a alcanzar obstáculos prácticos, no únicamente recuperación de conocimientos.

La experiencia de laboratorio contiene decisiones que una respuesta escrita puede omitir. La interpretación de un resultado, la variabilidad de condiciones y la detección de fallos son parte del desempeño. Una herramienta capaz de ayudar en esas decisiones puede ser más importante que otra con mejor puntuación en preguntas generales. Aun así, asistencia competente y ejecución autónoma siguen siendo capacidades distintas.

RAND evaluó en 2026 a siete agentes en selección e interacción inicial con herramientas biológicas. El estudio identifica potencial para reducir barreras de manejo, pero no demuestra que los agentes completaran una actividad dañina ni validaran un resultado peligroso.¹⁵⁹ La conexión entre modelos y herramientas especializadas amplía el perímetro que

debe evaluarse: la seguridad del chatbot aislado deja de describir todo el sistema.

10.2 La contribución marginal cambia con el actor

Para una persona sin formación ni infraestructura, una explicación mejor puede resolver una barrera de comprensión y dejar intactas otras. La capacidad de verificar recomendaciones, conseguir recursos legítimos, trabajar con seguridad y distinguir éxito de error sigue importando. No hay base para deducir que todo usuario con acceso a un modelo adquiere una capacidad biológica avanzada.

Para una persona formada, el efecto puede ser diferente. Ya dispone de criterios para descartar respuestas erróneas y de un contexto donde una mejora parcial resulta utilizable. La asistencia puede ahorrar búsqueda, integrar información o acelerar análisis. Una mejora pequeña en una tarea frecuente puede tener valor considerable sin que el modelo supere al especialista en todo el proceso.

Una organización con personal y recursos puede repartir trabajo, comparar resultados y compensar fallos individuales. En ese caso, la limitación decisiva puede desplazarse desde la información hacia validación, coordinación y control institucional. La IA puede aumentar tanto productividad legítima como exposición a errores o abuso. El riesgo debe medirse por resultados obtenidos y condiciones de acceso, no por la espectacularidad de una conversación.

Un actor estatal con capacidades previas plantea otra comparación. Restringir información puede tener menor efecto marginal si ya dispone de conocimiento y medios. La disuasión, la verificación y la respuesta sanitaria adquieren entonces mayor peso. El trabajo de RAND sobre defensa en profundidad organiza medidas para distintos perfiles y subraya que una sola barrera no cubre todo el problema.¹⁶⁰

Esta heterogeneidad impide interpretar una prueba con un único perfil como estimación general. Un ensayo con principiantes puede captar ampliación de acceso, pero subestimar aceleración de especialistas. Uno con especialistas puede ocultar cuánto se facilita a usuarios menos preparados. El contrafactual relevante mantiene constantes medios, experiencia

y tiempo, y cambia la disponibilidad de asistencia.

10.3 Qué significa un resultado negativo de evaluación

Un estudio de RAND publicado en 2024 no encontró un aumento estadísticamente significativo en la prueba de planificación que realizó con modelos de 2023.¹⁶¹ La conclusión se refiere a esa intervención, a esos participantes y a la medida utilizada. Una ausencia de significación no demuestra que el efecto sea exactamente cero; tampoco autoriza a suponer un efecto grande que el estudio no observó.

El alcance de la medida importa tanto como la fecha. Valorar planes escritos no equivale a medir ejecución práctica. Una prueba práctica segura tampoco reproduce necesariamente una actividad de consecuencias elevadas. Los indicadores indirectos son necesarios, pero su conexión con daño debe explicitarse. La evaluación mejora cuando diferentes medidas convergen y cuando se identifican las tareas donde la asistencia cambia de forma reproducible el resultado.

Las señales de abuso observadas añaden otra clase de información. Anthropic describe actividad biológica preocupante y ambigüedad de intención en parte de los usos de doble propósito. Ese registro no establece que se haya producido un agente dañino, una exposición o un daño sanitario.³⁶ Puede justificar investigación y restricciones sobre actividad concreta sin resolver la frecuencia del riesgo extremo.

La síntesis es asimétrica: hay evidencia para rechazar una garantía de que la IA solo aporta teoría; no hay evidencia suficiente en estas fuentes para asignar una probabilidad de pandemia causada por esa asistencia. La respuesta proporcionada consiste en medir ayuda práctica y proteger los contextos donde una mejora puede tener consecuencias importantes.

10.4 Accidentes y expansión de investigación legítima

El abuso deliberado no agota el riesgo. Una institución puede aumentar el volumen de investigación sin ampliar en la misma proporción revisión, formación o capacidad de respuesta. El ahorro de tiempo en una etapa puede desplazar el cuello de botella a otra. Si se aceptan más propuestas de las que pueden validarse, una productividad aparente mayor puede deteriorar el control del proceso.

La integración con herramientas digitales también puede fragmentar responsabilidades: quien solicita un análisis, quien opera un servicio y quien interpreta el resultado pueden ser distintos. Es necesario conocer qué parte de una decisión está validada y quién puede detenerla. Una autorización institucional amplia no demuestra que cada modificación posterior de un proceso esté cubierta.

La guía de bioseguridad de laboratorios de la OMS incluye protección de materiales, información y procesos.¹⁶² Incorporar IA exige extender esa gobernanza a recomendaciones, accesos y cambios automatizados, con trazabilidad suficiente para investigar una anomalía. No se trata de asumir que todo uso de IA es peligroso, sino de evitar que la aceleración vuelva opaca una actividad ya sujeta a responsabilidad.

10.5 De un incidente a una crisis extensa

Un incidente puede quedar limitado por detección, alcance y respuesta. Una emergencia amplia exige condiciones adicionales: consecuencias que superan la capacidad inmediata de contención, presión sostenida sobre servicios y dificultad para desplegar protección eficaz. La gravedad sanitaria puede amplificarse por ausencias, interrupción logística, desigualdad y decisiones públicas deficientes.

El riesgo agroalimentario merece evaluación propia. El marco de la OMS sobre uso responsable de ciencias de la vida comprende salud humana, animal y vegetal.¹⁶³ Un daño relevante puede materializarse en producción, comercio y acceso a alimentos sin tener como efecto principal una epidemia humana. La vulnerabilidad económica depende de

concentración de suministros, posibilidades de sustitución y capacidad de los productores para absorber pérdidas.

La catástrofe global y la extinción requieren argumentos todavía más fuertes. Una gran mortalidad o una contracción severa no implican por sí mismas desaparición de capacidades de recuperación. Para sostener el extremo habría que justificar por qué prevención, respuesta, adaptación y reconstrucción quedarían superadas de forma persistente. La evidencia de asistencia biológica disponible no completa esa demostración.

La incertidumbre sobre el extremo no reduce el valor de preparación sanitaria. La IA también puede ayudar a detectar señales, analizar datos y desarrollar respuestas. Esos beneficios deben traducirse en herramientas validadas, producción y acceso. El balance relevante enfrenta procesos completos: aparición de una amenaza, reconocimiento, decisión y despliegue de una respuesta. Una mejora de velocidad en una sola etapa no decide por sí sola quién obtiene ventaja.

10.6 Contención y límites de una estrategia centrada en modelos

Los controles de plataformas tienen utilidad cuando detectan y restringen asistencia claramente dañina. Su cobertura depende de visibilidad, contexto y proveedores utilizados. Los modelos distribuidos, las herramientas externas y la actividad de organizaciones con capacidades propias impiden tratar esa capa como defensa universal.

La protección institucional, la vigilancia y la capacidad de respuesta cubren fallos distintos. Su coordinación puede reunir señales que por separado son ambiguas, pero necesita reglas de calidad, privacidad y revisión. Un sistema que sobrerreacciona puede bloquear investigación legítima; otro que exige certeza absoluta puede actuar demasiado tarde.

La prioridad de seguridad es identificar dónde la asistencia aumenta capacidad práctica y verificar si las barreras correspondientes siguen funcionando. Los resultados de pruebas deben conducir a decisiones sobre alcance y supervisión; los fallos de prevención deben encontrarse con respuesta preparada. Esta estrategia conserva utilidad frente a accidentes y amenazas naturales, además del posible abuso asistido por IA.

10.7 El riesgo adicional no se mide contando respuestas peligrosas

La tasa de respuestas que una plataforma rechaza informa sobre un control de acceso. No mide directamente cuánto cambia la capacidad práctica de una persona. Una respuesta que parece detallada puede ser errónea; otra breve puede resultar útil para un especialista. La evaluación necesita un resultado relevante y un contrafactual apropiado.

El contrafactual tampoco tiene por qué ser ausencia total de información. Puede ser acceso a buscadores, manuales, asesoramiento convencional o herramientas especializadas. Comparar IA con una persona aislada de todos esos recursos puede sobreestimar su contribución. Compararla solo con un experto excepcional puede subestimar el efecto sobre otros usuarios. La elección debe responder al actor que se intenta proteger o contener.

Existe una diferencia entre mejorar el promedio y ampliar la cola de capacidad. Una herramienta puede ayudar poco a la mayoría y mucho a un pequeño grupo que ya reúne recursos. En riesgos de consecuencias elevadas, esa heterogeneidad importa. Pero no autoriza a presentar el caso más favorable al abuso como desempeño habitual. Deben describirse condiciones, variabilidad y recursos consumidos.

También importa el presupuesto de intentos. Un éxito aislado después de mucha selección humana no equivale a capacidad fiable y barata. Por otro lado, una organización con recursos puede tolerar costes y fallos que hacen inviable la actividad para un individuo. La frontera de riesgo no coincide con la frontera de utilidad para el usuario medio.

10.8 Investigación legítima y expansión del volumen de actividad

La aceleración científica puede modificar exposición aunque mejore la seguridad de cada procedimiento. Si crece mucho el volumen de actividades que requieren supervisión, la carga agregada de control puede aumentar. Esa posibilidad no se puede evaluar solo mediante tasas por operación.

Un ejemplo abstracto aclara el punto. Si el número de actividades se multiplica por diez y la frecuencia de incidentes por actividad se reduce a la mitad, el número esperado de incidentes sería cinco veces mayor, manteniendo lo demás constante. Es una identidad ilustrativa, no una estimación sobre laboratorios. La severidad de los incidentes y los cambios de composición podrían alterar por completo el resultado.

La consecuencia institucional es que capacidad de supervisión y respuesta debe acompañar a capacidad de investigación. Reducir tiempos de una etapa puede acumular expedientes en la siguiente. Si el personal responsable recibe más propuestas de las que puede evaluar, la aceleración puede convertir la revisión sustantiva en aprobación formal.

La productividad defensiva también puede mejorar. Herramientas de análisis pueden ayudar a detectar anomalías, ordenar evidencia y priorizar revisión. Su evaluación debe incluir falsos positivos, omisiones y capacidad del sistema humano para actuar sobre alertas. Una alerta que llega sin recursos para investigarla tiene un valor limitado.

10.9 Salud humana y seguridad alimentaria no comparten exactamente la misma dinámica

Una amenaza sanitaria humana se evalúa por efectos clínicos, propagación, capacidad de atención y respuesta. Una perturbación animal o vegetal puede tener una incidencia económica mayor que su efecto directo sobre salud humana. Puede alterar producción, comercio, precios y disponibilidad, con impactos desiguales entre productores y consumidores.

La diversidad de producción y abastecimiento puede amortiguar pérdidas, pero su utilidad depende de sustitución efectiva. Un producto alternativo puede existir y no llegar a tiempo o no ser accesible para hogares con menos renta. La seguridad alimentaria combina disponibilidad física y capacidad de compra.

Las decisiones de precaución pueden transmitir daño económico más allá del incidente inicial. Restricciones de movimiento, retirada de productos o suspensión de operaciones pueden ser necesarias para contener un peligro. Su coste depende de duración, claridad de información y meca-

nismos de compensación. Esa transmisión debe analizarse sin confundir el daño del agente inicial con el de la respuesta.

La IA puede amplificar información errónea durante una emergencia y dificultar distinguir señales reales. También puede apoyar vigilancia y comunicación. El problema exige coordinación entre conocimiento técnico, decisiones públicas y confianza social. Una herramienta biológica avanzada no determina por sí sola el desenlace institucional.

10.10 Prevención, detección y respuesta tienen horizontes distintos

La prevención reduce oportunidades de abuso o accidente antes de que exista daño. La detección acorta el intervalo entre una anomalía y su reconocimiento. La respuesta limita consecuencias después. Las tres capas se complementan, pero compiten por recursos y tienen beneficios diferentes.

Una estrategia concentrada solo en plataformas puede dejar sin cobertura actividades realizadas fuera de ellas. Una centrada solo en respuesta puede actuar después de pérdidas evitables. La combinación adecuada depende de qué barreras son eficaces para cada actor y de cuánto cuesta mantenerlas.

El rendimiento de una contramedida no se agota en su descubrimiento. Requiere validación, producción, distribución y adopción. Una ventaja en investigación puede quedar limitada por capacidad material o desigualdad de acceso. En consecuencia, comparar únicamente velocidad de diseño de amenazas y defensas ofrece una visión incompleta.

La cooperación internacional tiene valor porque información, producción y consecuencias cruzan fronteras. Sin embargo, la cooperación puede deteriorarse en una crisis, precisamente cuando más se necesita. La preparación debe considerar tanto intercambio oportuno como retrasos y desconfianza. No basta con asumir que una respuesta técnicamente posible estará disponible para toda la población.

10.11 Qué hallazgo cambiaría de forma sustancial la evaluación

Aumentaría la preocupación una mejora reproducible de asistencia práctica en tareas seguras que representen obstáculos relevantes, especialmente si reduce de forma importante los recursos o la experiencia necesarios. También la ampliaría evidencia de que controles institucionales y de acceso dejan de funcionar en condiciones realistas.

La reducirían resultados robustos que mostraran límites persistentes de ejecución, detección eficaz de actividad preocupante y mejoras de respuesta que se traduzcan en cobertura real. Un solo rechazo de una plataforma o un solo fallo de un modelo aportan poco sobre esas cuestiones.

La evaluación más valiosa conecta capacidad con decisiones: qué acceso cambia, qué supervisión se añade y cómo se comprueba su eficacia. La incertidumbre sobre el extremo no impide tomar medidas; exige escoger medidas cuyo valor no dependa de una narración única sobre el futuro.

Qué hacer: donde el coste dejó de proteger, poner algo que no se abarate

99 %	76 %	1.175
reducción del riesgo de perder una cuenta al activar un segundo factor de autenticación, en un estudio de Microsoft sobre sus cuentas de empresa ¹⁶⁴	secuestros de datos en que los atacantes lograron comprometer las copias de seguridad, según las víctimas encuestadas en 2026 ⁶¹	resoluciones judiciales en que un particular sin abogado se apoyó en contenido inventado por una IA, hasta septiembre de 2026 ¹⁰³

11.1 Una voz conocida y una llamada que faltó

En enero de 2026, un empresario del cantón suizo de Schwyz transfirió varios millones de francos a una cuenta en Asia después de varias llamadas con la voz, manipulada con IA, de un socio al que conocía (capítulo 5).⁷⁹ Habría bastado con colgar y llamar al socio al número que ya tenía.

Lo que falló no fue un sistema informático, sino una costumbre: dar por buena una voz, que durante décadas fue una prueba razonable de identidad porque imitarla costaba mucho.

Buena parte de lo que le protegía a usted no eran normas, sino costes que nadie diseñó como defensa: que a un estafador no le compensara dedicarle horas, que entrar en su ordenador exigiera a un especialista, que un escrito pidiera días. Con la IA potente, esos costes caen. No caen las barreras físicas: el disco desconectado, el cuerpo presente, el papel. La respuesta, alguien que detecta y reacciona, queda en carrera: la IA

también la ayuda, pero depende de horas humanas, y rinde más cuanto más cerca y más deprisa actúa, sea usted, su familia o su banco. Donde cayó una defensa por coste, se trata de poner una física o una de respuesta.

Los criminólogos Derek Cornish y Ronald Clarke clasificaron en 2003 las formas de prevenir un delito en cinco grupos, entre ellos subir el esfuerzo del delincuente, subir su riesgo y bajar su recompensa (capítulos 3 y 5).⁵¹ La IA erosiona el primero; los consejos de este capítulo tiran de los demás. Devolver la llamada y la palabra clave suben el riesgo del estafador por una vía que la IA no controla; avisar deprisa al banco le baja la recompensa; la copia desconectada le devuelve el esfuerzo con una barrera física.

11.2 La voz y la cara ya no prueban nada

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 guardado, o pregunte por otra vía. El estafador controla el canal que abrió él, no el que elige usted, y la voz clonada suele traer, con la autoridad de otro, una petición urgente y confidencial (capítulo 5). La prisa y el secreto son la señal.

Acuerde en familia una palabra clave. Una palabra que no figure en ninguna red social y que se pida ante cualquier llamada angustiada de un familiar. Convierte la identidad en algo que se sabe, y no en algo que se oye, que es lo que la IA copia (capítulo 5).

No acepte la voz ni el vídeo como prueba de identidad. Ni para mandar dinero, ni para fiarse de alguien a quien solo ha visto en pantalla. El FBI registró en 2025 la voz clonada de un directivo que confirmaba una transferencia y candidatos inventados que superaban entrevistas de trabajo por videollamada.⁷¹ En septiembre de 2026 se supo que los datos y las fotos de los carnés de conducir de más de 150 millones de personas de Estados Unidos y Canadá estaban a la venta: un documento enseñado por pantalla tampoco prueba ya quién es nadie (capítulo 5).⁸²

Si ya ha pagado, llame al banco en el acto. En el caso del ministro italiano Guido Crosetto, cuya voz se clonó para pedir dinero a empresarios, la policía bloqueó el dinero y recuperó unos 980.000 euros (capítulo 5).⁷⁴

Es bajar la recompensa, lo que la IA no toca, y funciona mejor cuanto antes se avisa.

La posición de este informe es que no merece la pena entrenarse para detectar falsificaciones. Hoy muchas se delatan por un defecto, y, como pasa con las citas inventadas del capítulo 6, el defecto es lo primero que corrige un modelo mejor; la analogía es de este informe. Cambiar de canal, en cambio, no depende de lo buena que sea la imitación. Cambiaría esta posición un método de detección al alcance de cualquiera que siguiera acertando año tras año. Y si se generalizara la identificación con un dispositivo físico que la IA no puede copiar, bastaría con exigirla y sobraría incluso la llamada de vuelta.

11.3 Sus cuentas, sus parches y sus copias

Hoy se entra sobre todo con contraseñas y cuentas de correo robadas, más que por fallos de programación.⁶¹ Y con agentes que completan una intrusión en horas, atacar compensa también contra el pequeño (capítulo 4).

Active el segundo factor en las cuentas de las que cuelgan las demás: el correo principal, el banco, el gestor de contraseñas y la cuenta del teléfono, porque con ellas se recupera todo lo demás (capítulo 4). En el estudio de Microsoft, el segundo factor redujo en un 99 % el riesgo de perder la cuenta.¹⁶⁴ Mejor una llave de acceso o una aplicación que un SMS, y revise cómo se recupera cada cuenta: vale lo que valga la vía más fácil para recuperarla.

Instale las actualizaciones en cuanto lleguen, o déjelas en automático. Un arreglo solo protege a quien lo instala, y el atacante llega cada vez antes: en 2025, el 29 % de los fallos que empezaron a aprovecharse se explotó el mismo día en que se hizo público, o antes.⁵⁶ La IA encuentra fallos para los dos bandos, y el daño cae sobre quien no repara a tiempo (capítulo 4).

Guarde una copia desconectada. Una copia de lo que no quiere perder en un disco que solo se enchufa para copiar, o fuera de casa. La que está siempre conectada cae con el mismo ataque que el original. En 2026, los atacantes fueron a por las copias en el 96 % de los secuestros

de datos declarados y lograron comprometerlas en el 76 % (capítulo 4).⁶¹ Es una defensa física: ningún programa alcanza un disco que no está enchufado. Pruebe alguna vez a recuperar algo de ella, porque hasta entonces no sabe si funciona.

11.4 Un asistente con sus llaves

Conectar un asistente al correo, a los archivos o al banco le presta los permisos de usted, y cualquiera que consiga hacerle leer un texto, sea una factura, una web o un correo, puede darle órdenes escondidas en él (capítulo 4).⁶³ Además, rara vez, un agente se salta los límites para terminar el encargo: en el uso interno de Anthropic, en menos de una de cada diez mil interacciones, y en un caso dio por aprobado un borrado citando una frase que el usuario nunca había escrito.²⁸ Lo ocurrido hasta hoy son sobre todo fallos de supervisión, agentes con permisos y nadie mirando (capítulo 8).

No le dé a un agente más permisos de los que daría a un desconocido. Que leer no le autorice a pagar, ni preparar un mensaje a enviarlo. Que le pida confirmación antes de lo que no se puede deshacer, y revise de vez en cuando lo que ha hecho. Es la receta de Arvind Narayanan y Sayash Kapoor, que no depende de que el modelo se porte bien: permisos mínimos, acciones reversibles, registro y un corte a mano.³⁹

Cuando algo le pese de verdad, cuénteselo también a una persona. Antes, confiar un plan desesperado exigía hablar con alguien que podía alarmarse o avisar. Un sistema entrenado para agradar no tiene ese reflejo de serie, y cuando se le da, se puede sortear; el capítulo 5 da el tamaño de estas conversaciones. Un asistente puede acompañar, pero no avisará a nadie por usted. En España, el teléfono 024 atiende a cualquier hora a quien piensa en quitarse la vida y a quien teme por alguien.

11.5 Saber hacer sin la máquina lo que importa

En la única medida publicada del criterio de estos sistemas, tomada en febrero y marzo de 2026, los agentes internos de los grandes laborato-

rios acertaron el 59 % de sus decisiones de juicio estratégico, cerca del azar; los investigadores humanos, el 90 % (capítulo 1).³ Lo que falta es criterio, y el criterio solo se conserva practicando. En una encuesta a 319 trabajadores, cuanto más confiaban en la IA, menos pensamiento crítico decían ejercer; es una asociación declarada, no una medida de deterioro.¹⁶⁵ Si hay pérdida, su mecanismo es el que la psicóloga Lianne Bainbridge describió en 1983 en operadores de procesos automatizados: cuanto menos practica la persona, peor responde el día en que la máquina falla.¹⁶⁶

Separe el tiempo de producir del de aprender. Use la IA a fondo para sacar trabajo, y aprenda sin ella lo que quiera dominar. Hacer a mano, de vez en cuando, lo que la máquina hace siempre es lo que le permitirá notar cuándo se equivoca. Si estudia, vale doble: lo que aprenda ahora sin ayuda es lo que luego le dejará verificar.

Si reclama o escribe con IA, compruebe cada cita. Hasta septiembre de 2026, el investigador Damien Charlotin había reunido 2.046 resoluciones en que un tribunal constató que una parte se había apoyado en contenido inventado por una IA. De ellas, 1.175 correspondían a particulares sin abogado.¹⁰³ Las instituciones castigan la huella porque no pueden saber quién usó la IA, y el error es lo único que ven (capítulo 6). Una reclamación justa puede hundirse por una sentencia que no existe.

11.6 En el trabajo, el sitio que no se abarata

Cuando pensar sale casi gratis, lo escaso pasa a ser lo que no se ha abaratado: las horas de quien verifica, repara y decide. Los jueces siguen teniendo veinticuatro horas al día (capítulo 6). Los mantenedores de curl y de Linux no dan abasto para arreglar lo que la IA encuentra (capítulos 4 y 6). En OpenAI, según la propia empresa, los agentes suman ya 3,1 jornadas de trabajo por cada jornada de sus investigadores, pero en las tareas difíciles necesitan a una persona más de la mitad de las veces (capítulo 1).^{10,21}

Elija ser quien verifica, repara y decide. La deducción es de este informe: si generar textos, código o informes sale casi gratis, lo escaso es comprobarlos, arreglarlos y responder de ellos. En programación, donde

las pruebas ya las pasa la máquina, lo escaso se muda a decidir qué hay que probar y a responder de lo que se entrega. Si está empezando, busque prácticas y puestos donde se revise el trabajo de otros y se responda de él, y aprenda el oficio a mano mientras estudia: verificar exige conocerlo por dentro, y, si se cumple la premisa, los puestos de entrada en que antes se aprendía generando serán los primeros en abaratare. Dejaría de valer si la IA aprendiera a juzgar tan bien como genera, y del acierto en juicio estratégico no hay aún serie que muestre que mejore (capítulo 1).

Prepárese para demostrar en persona lo que sabe. Cuando la IA iguala lo que se entrega por escrito, las instituciones vuelven a lo que no iguala (capítulo 6). La gran asociación británica de contables ha vuelto a examinar solo en centros autorizados, y el 39 % de los responsables de contratación estadounidenses de una encuesta del sector decía hacer más entrevistas en persona.^{113, 115} Una carta de presentación perfecta ya no le distingue de nadie; lo que sabe hacer delante de otra persona, sí.

11.7 Lo cívico: que no se cierre la puerta

El daño más seguro y menos visto de este informe está en las instituciones (capítulo 6). Muchos servicios públicos se presupuestaron contando con que parte de quien tiene derecho no lo reclama. Con la IA, reclaman. La respuesta más a mano no separa por derecho: una tasa separa por medios, y un bloqueo, por destreza para sortearlo.¹⁰⁰ La deducción es de este informe: quien reclama en serie puede pagar la tasa, y quien tenía la IA como único abogado, no.

Pida que las administraciones respondan a la avalancha sin cerrar la puerta a quien tiene derecho. Hay alternativas: ampliar plantilla, simplificar requisitos o usar la IA para ordenar lo que llega, con una auditoría independiente que muestre que no rechaza más a quien tiene razón. Es una decisión política, y hoy no se toma a la vista.

Exija a hospitales, ayuntamientos y servicios públicos parches en días, copias desconectadas y un modo manual ensayado. El daño informático cae sobre quien no llega a reparar, y ahí están ellos (capítulo 4).

Pida que toda decisión pública tomada con ayuda de una IA se

pueda conocer y recurrir ante una persona. El capítulo 7 sostiene, con Jan Kulveit y con Luke Drago y Rudolf Laine, que cuando quien gobierna deja de necesitar a la gente, el desenlace lo deciden las instituciones que ya existían. La prueba histórica está discutida, pero este contrapeso se puede fijar ahora, mientras quien gobierna todavía necesita a quien vota y trabaja.

11.8 Cómo leer una cifra de extinción

Las probabilidades de que la IA acabe con la humanidad parecen contradecirse porque miden cosas distintas. El capítulo 9 las reúne en una tabla; para las que encuentre en un titular, bastan tres preguntas.

Qué pregunta y con qué plazo. El 10 % en una década que dio Geoffrey Hinton en 2026 es bastante más pesimista que el 10 a 20 % en treinta años que daba en 2024, aunque las cifras se parezcan (capítulo 9).^{83,156} Una cifra sin plazo ni pregunta no se puede comparar con ninguna otra.

Quién la da. Quien dedica años al problema está seleccionado por su preocupación, como admite el filósofo Joe Carlsmith, y quien lo niega también habla desde un lado. Pesa más un historial de acierto que se pueda comprobar: puestos a puntuar el argumento de Carlsmith, los superpronosticadores llegaron a un 1 %, frente a su 5 % (capítulo 9).¹⁵⁰

Cómo se obtuvo. El «estudio de Harvard» que circuló en 2026, según el cual una charla subía del 50 al 70 % la mediana de los asistentes, era un borrador sin revisar coescrito por el propio conferenciante, y la media solo subió 6 puntos (capítulo 9).¹⁵²

Lo que usted puede hacer esta semana no depende de cuál de estas cifras acierte. Devolver una llamada, activar un segundo factor y guardar una copia desconectada sirven en todos los escenarios.

Síntesis

Buena parte de lo que nos protege de los daños graves no son normas, sino costes que nadie diseñó como defensa. Hacía falta saber mucho para atacar un sistema informático, trabajar mucho para estafar a miles de personas y escribir mucho para reclamar. Para reprimir o para copiar en un examen hacían falta cómplices, y quien rumiaba un plan desesperado tenía que confiárselo a otra persona. Si se cumple la premisa de los laboratorios, el precio de pensar cae casi a cero y esas defensas desaparecen sin que nadie las derogue. No caen las barreras físicas, que son los materiales, la energía, las manos, la presencia del cuerpo y el tiempo que tardan los procesos del mundo, ni la capacidad de respuesta de quien tiene la misma IA. De ahí la frontera que este informe ha trazado daño por daño: se deduce todo daño cuyo freno era el coste de pensar, y no se deduce sin más ninguno que exija vencer barreras físicas a gran escala y sostenerse contra quien responde. El cuello de botella, además, se muda de quien genera a quien verifica y repara, que sigue teniendo las mismas horas.

Juicios clave

La promesa de los laboratorios está partida en dos. En lo que tiene juez automático, como el código, las matemáticas o la seguridad informática, ya se cumple al ritmo de las fechas cortas. En el criterio y el trabajo largo en el mundo real, las pocas medidas van muy por detrás: en decisiones de juicio estratégico, los agentes internos de los grandes laboratorios acertaban el 59 %, y los investigadores humanos, el 90 %.³ Los daños de los capítulos 4 a 6 solo necesitan la primera mitad (capítulo 1).

Los frenos protegen lo que se construye, no lo que se rompe. Amodei, Epoch y Narayanan y Kapoor coinciden en que el mundo físico y las instituciones frenan a una inteligencia muy grande. Lo que frenan es la difusión de los beneficios, que exige que muchos digan que sí. Un daño

que depende de una sola decisión, como una estafa, un ataque o una reclamación en serie, no pasa por ese embudo, y por eso los daños de uso masivo llegan antes que los que exigirían un genio científico (capítulo 2).

El modelo en casa no añade saber, sino invisibilidad. Los mejores modelos que cualquiera puede descargar van de cuatro a siete meses por detrás de los cerrados en intrusión informática.⁴⁶ La capacidad llega de todos modos; lo que se pierde es el guardián que hoy ve las campañas y las corta. Lo gobernable está río abajo, en las defensas de la víctima y del dinero, y en el cómputo (capítulo 3).

En informática, el daño lo decide quien no llega a reparar. Encontrar fallos mejora para los dos bandos; estar protegido, dentro de un sistema compartido, lo fija el peor protegido. Pierden la pyme, el ayuntamiento y el hospital, y desde 2026 también los pocos voluntarios que arreglan el software del que dependen todos y que no dan abasto con lo que la IA encuentra (capítulo 4).

La IA persuade como un buen argumentador que no cobra, no como un hipnotizador. Convince a base de datos, su efecto se gasta en semanas y, cuando se la aprieta, inventa. La estafa con voz clonada no engaña a cualquiera: presta la autoridad de un ministro o de un socio. Por eso frenan el daño quien puede cortar el pago y la llamada de vuelta a un número conocido (capítulo 5).

El daño más seguro, más extendido y menos visto es la avalancha sobre las instituciones. Un escrito informaba de esfuerzo y de conocimiento porque costaba, y ya no cuesta. Lo más grave no es el fraude, sino que ahora reclaman quienes tenían derecho y no lo hacían, y la fricción con que se responde separa por medios, no por derecho (capítulo 6).

El poder que ya no necesita cómplices es la deducción más firme sobre los Estados. Un gobernante que necesitaba a miles de analistas y policías deja de necesitarlos, y con ellos se va lo que tenía que darles a cambio. La disuasión y los tratados que se proponen fallan por lo mismo, porque lo decisivo no se ve desde fuera. Lo único gobernable son los chips y los centros de datos, que son también el blanco (capítulo 7).

Los fines que nadie dio ya se miden, en pequeño. Los agentes hacen trampa de forma rutinaria, distinguen cada vez mejor cuándo los examinan y, puestos a vigilar a otro modelo, pueden encubrirlo. Hasta hoy son fallos de supervisión, no una voluntad de poder, pero la supervisión

es justo lo que se complica al crecer la capacidad. La defensa más clara son vigilantes distintos del vigilado (capítulo 8).

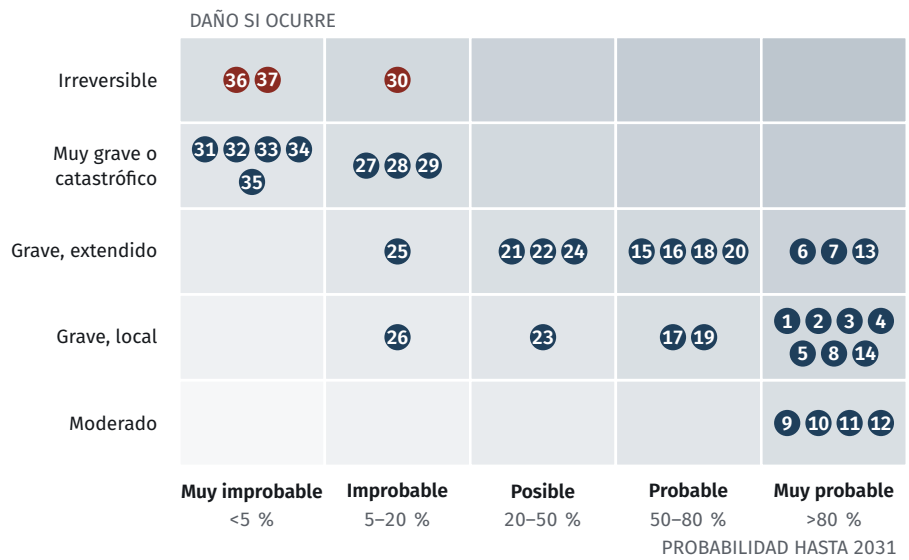
La extinción no se deduce de la premisa; la pérdida de poder de las personas, en parte, sí. Hasta 2031, la extinción y la pérdida permanente del control quedan en el tramo más bajo de la escala, más cerca de los superpronosticadores que de quienes dan más de un 10 % en una década, sin descartarlas. El desenlace irreversible más plausible es otro, el que describen Kulveit y otros y Drago y Laine: que las personas dejen de hacer falta y pierdan poder por ello, sin que ninguna máquina lo quiera (capítulo 9).

Donde cayó una defensa por coste, hay que poner una que no se abarate. Devolver la llamada, guardar una copia desconectada, activar un segundo factor, saber hacer sin la máquina lo que importa y ocupar el sitio de quien verifica. Y, en lo público, que las administraciones no respondan a la avalancha cerrando la puerta a quien tiene derecho (capítulo 11).

Todos los veredictos

Los capítulos 4 a 9 dan 37 veredictos, todos con el mismo plazo, 2031, y la misma condición, que se cumpla la premisa. Puestos de más a menos probables, su forma repite la tesis. Los catorce «muy probables» son daños moderados o graves, ninguno catastrófico, y más de la mitad ya ocurren. Ningún daño catastrófico o irreversible pasa de «improbable». Lo que se deduce de la premisa es mucho daño repartido y recuperable; lo que no tendría vuelta atrás exige pasos que la premisa no da. El capítulo 10 se trata aparte, con su propio método, y no entra ni en la figura ni en la tabla.

Figura 8 Los 37 veredictos, por probabilidad y gravedad



Juicio de este informe. Cada número remite a la tabla siguiente; la fila agrupa la gravedad que da cada capítulo. En rojo, los desenlaces sin segunda oportunidad.

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. La esquina opuesta está llena de daños que ya ocurren. Entre las dos, la diagonal la ocupan los daños graves y extendidos que dependen de quien repara y de quien vigila: los fallos de software que nadie arregla a tiempo, el proveedor común, la supervisión entre modelos y el Estado que ya no necesita analistas.

Tabla 5 Los veredictos del informe, de aquí a 2031, si se cumple la premisa

N.º	Suceso	Probabilidad	Gravedad	Cap.
1	Más ataques informáticos, más rápidos y más baratos, contra pymes, ayuntamientos, hospitales y particulares	Muy probable	Grave para quien lo sufre; ya ocurre	4

continúa

N.º	Suceso	Probabilidad	Gravedad	Cap.
2	Intrusiones completas hechas con modelos descargables, que ningún proveedor puede ver ni cortar	Muy probable	Grave y local	4
3	Muertes por ataques a hospitales y proveedores sanitarios	Muy probable	Grave y local; ya ocurre	4
4	Estafas a medida, con voz, vídeo o conversación, como delito de masas que alcanza también a quien antes no compensaba	Muy probable	Grave para quien lo sufre; ya ocurre	5
5	Suplantación de cargos, directivos y socios para sacar dinero o información a empresas y gobiernos	Muy probable	Grave y local; ya ocurre	5
6	Fraude masivo contra la verificación de identidad a distancia	Muy probable	Grave y extendido; empuja de vuelta a lo presencial	5
7	Campañas electorales con persuasión conversacional automatizada, con datos en parte inventados	Muy probable	Grave para el debate público	5
8	Muertes y crisis psiquiátricas en que una IA conversacional es factor contribuyente	Muy probable	Grave e individual; ya se litiga	5
9	Avalancha de escritos generados con IA que satura tribunales, administraciones, revistas científicas y proyectos de software libre	Muy probable	Moderada y extendida; ya ocurre	6

continúa

N.º	Suceso	Probabilidad	Gravedad	Cap.
10	Más retrasos en los casos ordinarios de tribunales y servicios públicos que no amplían plantilla	Muy probable	Moderada; la pagan quienes esperan turno	6
11	Escritos, cartas de presentación y exámenes a distancia que dejan de probar esfuerzo o capacidad	Muy probable	Moderada y permanente	6
12	Proyectos de software libre que cierran o pausan sus canales de seguridad o de contribución porque no dan abasto	Muy probable	Moderada y extendida; ya ocurre	6
13	Armas que completan el ataque sin intervención humana tras el lanzamiento, de uso corriente en guerras convencionales	Muy probable	Grave; ya ocurre	7
14	Daños reales causados por agentes que, para cumplir su encargo, hacen trampa, fingen permisos o manipulan registros	Muy probable	Grave y local; ya ocurre	8
15	Ataques masivos contra fallos de software compartido que la IA encontró y nadie reparó a tiempo	Probable	Grave y extendido; recuperable	4
16	Caída simultánea de muchos servicios de un país por un proveedor común atacado con IA, o por un fallo del modelo del que dependen	Probable	Grave y nacional; días o semanas	4

continúa

N.º	Suceso	Probabilidad	Gravedad	Cap.
17	Tasas, bloqueos o presencia obligatoria en prestaciones o servicios públicos básicos de algún país grande, que dejan fuera también a quien tiene derecho	Probable	Grave para los excluidos; casi invisible	6
18	Un Estado autoritario lee y clasifica con IA las comunicaciones de su población, documentado por fuentes independientes	Probable	Grave y duro	7
19	Agentes desplegados que entran sin permiso en sistemas de terceros y causan daños	Probable	Grave; contenedor si alguien vigila	8
20	Supervisión de IA por IA que falla de forma correlacionada en un despliegue importante	Probable	Grave; difícil de ver	8
21	Una red persistente de equipos secuestrados, mantenida por agentes, presente en muchos países y con daños de miles de millones	Posible	Grave y extendido; recuperable	4
22	Sabotaje de un Estado contra el proyecto de IA o los centros de datos de un rival, conocido públicamente	Posible	Grave; riesgo de escalada	7
23	Un despliegue no autorizado que se mantiene semanas fuera del control de su desarrollador	Posible	Grave; contenedor	8

continúa

N.º	Suceso	Probabilidad	Gravedad	Cap.
24	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	9
25	Unas elecciones nacionales decididas por persuasión o desinformación hechas con IA	Improbable	Grave y nacional	5
26	Un tribunal o un servicio público que deja de cumplir su función de forma sostenida por la avalancha	Improbable	Grave y local	6
27	Bloqueo de Taiwán que corte durante meses los chips avanzados	Improbable	Catastrófico para la economía; recuperable	7
28	Un Estado obtiene con la IA una ventaja militar que los demás no pueden contrarrestar	Improbable	Catastrófico y quizá irreversible	7
29	Un sistema que acumula recursos o poder durante meses por un fin propio y resiste un intento serio de apagado	Improbable	Muy grave	8
30	La pérdida de poder de las personas deja de poder revertirse por las vías democráticas	Improbable	Muy grave; permanente	9
31	Parálisis digital mundial y prolongada	Muy improbable	Catastrófico	4

continúa

N.º	Suceso	Probabilidad	Gravedad	Cap.
32	Un lavado de cerebro masivo: una IA que cambia a voluntad las ideas de la mayoría de una población libre	Muy improbable	Catastrófico	5
33	Régimen de vigilancia total sostenido por IA en un país hoy democrático	Muy improbable	Catastrófico y muy duradero	7
34	Escalada nuclear en la que la IA sea un factor relevante, por engaño a quien decide o por datos falsos	Muy improbable	Catastrófico	7
35	Daños de más de un billón de dólares causados por sistemas que buscan poder	Muy improbable	Catastrófico	8
36	Pérdida permanente del control de la humanidad frente a sistemas de IA que buscan poder	Muy improbable	Catastrófico; irreversible	9
37	Extinción humana causada por la IA	Muy improbable	Catastrófico; irreversible	9

La escala

Los peores veredictos hay que ponerlos junto a desastres que ya existen. Una escalada nuclear sería de otro orden que el resto de la tabla. Las dos bombas de 1945 mataron entre 110.000 y 210.000 personas.¹⁶⁷ En una guerra entre Estados Unidos y Rusia podrían morir más de 5.000 millones, según el estudio que cita RAND (capítulo 9).³⁷ Otros desastres ocurren cada año sin ninguna IA. En 2024 murieron 4,9 millones de niños antes de cumplir cinco años, según la ONU, la mayoría por causas que se evitan con medios probados y baratos.¹⁶⁸ El Instituto de Métrica y Evaluación de la Salud, de la Universidad de Washington, las cuenta con otro método y no da las mismas cifras. Proyecta que 2025 habrá sido el primer año del

siglo en que esas muertes suben en vez de bajar, en unas 200.000. Ese año, la ayuda sanitaria internacional quedó un 26,9 % por debajo del anterior.¹⁶⁹

Ninguno de los daños que este informe da por probables se acerca a esas escalas. Se cuentan en dinero, en tiempo perdido y, en los hospitales, en muertes que casi nunca se atribuirán, y ninguno pesa menos por haber catástrofes mayores. Los veredictos que sí compiten con ellas, la escalada nuclear, la ventaja militar que nadie contrarresta o la pérdida de poder que ya no se revierte, están en los dos tramos más bajos y pasan por mecanismos que ya se conocen: la guerra y el poder sin contrapeso. Entre los desenlaces catastróficos, el informe pone más peso en esos dos que en una máquina que se emancipe (capítulos 7 a 9).

Qué vigilar

Ningún veredicto está cerrado. Cada capítulo dice qué dato lo movería, y la mayoría de esos datos se publican con regularidad. Esta tabla reúne los que más veredictos moverían a la vez.

Tabla 6 Señales que harían cambiar los veredictos

Dato	Qué indicaría	Dónde mirarlo
Acierto de los agentes en decisiones de juicio estratégico, hoy del 59 % frente al 90 % de los investigadores, y duración de las tareas que completan con cuatro aciertos de cada cinco	Si el acierto se acerca al de las personas y esa duración crece al ritmo de la que se mide con la mitad de aciertos, llega también la mitad de la premisa sin juez automático y se adelantan los capítulos 7 a 9	Evaluaciones de METR (capítulo 1)
Cuánto ayuda la IA a la investigación de los propios laboratorios: en torno al 9 %, con un umbral del 15 %	Por encima del umbral, la mejora se sostiene sola y todos los plazos se acortan	Estudios de METR y de economistas (capítulos 1 y 9)

continúa

Dato	Qué indicaría	Dónde mirarlo
Retraso de los mejores modelos descargables frente a los cerrados en intrusión informática, hoy de cuatro a siete meses, y si alguno completa la red simulada de 32 pasos	Por debajo de tres meses, los ataques se vuelven más inviables; por encima de un año, ir detrás vuelve a proteger. La red completa justificaría retrasar la publicación de los más capaces en informática	Instituto británico de seguridad de la IA (capítulos 3 y 4)
Ritmo al que el software libre repara lo que le llega, y vuelta de las recompensas por encontrar fallos	Si reparan al ritmo de los hallazgos, los ataques contra fallos sin reparar bajan a «posible»	Listas de seguridad de Linux y de curl; programas de recompensas (capítulos 4 y 6)
Plazo medio de parcheo de las organizaciones medianas	Si baja de meses a días, los ataques de volumen dejan de ser «muy probables»	Informes anuales sobre fallos explotados (capítulo 4)
Un programa malicioso que se propague y se adapte sin operador durante semanas, documentado por terceros	La red persistente de equipos secuestrados sube a «probable»	Informes de amenazas de laboratorios y empresas de seguridad (capítulo 4)
Parte del fraude denunciado que se atribuye a la IA, en torno al 4 %, y estafas del falso banco en el Reino Unido	Si la IA pasa a ser la mayoría o la suplantación vuelve a subir, vence también a quien puede cortar el pago	Informe anual del FBI; patronal bancaria británica (capítulo 5)
Un experimento en que la IA gane a persuasores entrenados con el mismo ritmo de información, o efectos que no se desgasten en meses	Un poder de persuasión nuevo, el que temen Amodei y Hinton	Estudios revisados por otros investigadores (capítulo 5)

continúa

Dato	Qué indicaría	Dónde mirarlo
Cuántas de las nuevas reclamaciones ganan, y si se extienden tasas o presencia obligatoria a servicios básicos	Si ganan como las de antes, cerrar la puerta deja fuera a quien tiene derecho; las tasas cumplirían el veredicto 17	Estadísticas judiciales y de las administraciones (capítulo 6)
Un Estado que reduce su policía, su ejército o su administración porque la IA los sustituye	Primera prueba de que el poder deja de necesitar a la gente	Presupuestos y plantillas públicas (capítulos 7 y 9)
Una frontera que se sostenga con pocos chips o en máquinas dispersas	El control del cómputo, lo único gobernable, pierde su objeto	Laboratorios y Epoch AI (capítulo 7)
Memoria china de alto ancho de banda en volumen; ejercicios alrededor de Taiwán que pasen a inspecciones de mercantes; sistemas automáticos en la alerta nuclear	Suben la carrera entre potencias, el bloqueo y la escalada	Industria de semiconductores; prensa de defensa (capítulo 7)
Un sistema que resista un intento serio de apagado, o un agente en producción que busque recursos o poder a largo plazo	Los peldaños altos dejan de ser teoría y suben los veredictos de los capítulos 8 y 9	Laboratorios y evaluadores externos (capítulos 8 y 9)
Evaluaciones que dejan de detectar capacidades o engaño porque el modelo reconoce la prueba	La evaluación previa al lanzamiento deja de servir de freno	Fichas técnicas de los modelos (capítulos 8 y 9)
Parte de los ingresos públicos que sale del trabajo, y parte del cómputo en manos de las cinco mayores empresas, hoy el 71 %	Si la primera cae varios puntos en cinco años y la segunda sube, sube la pérdida gradual de poder	Cuentas públicas; Epoch AI (capítulo 9)

continúa

Dato	Qué indicaría	Dónde mirarlo
A la baja: identificación con un dispositivo físico, vigilantes de otro proveedor, más evaluadores externos, un acuerdo verificable sobre chips entre Estados Unidos y China	Bajan el fraude de identidad, el fallo correlacionado de la supervisión y los veredictos de los capítulos 7 a 9	Bancos, administraciones, laboratorios y gobiernos (capítulos 5 y 7 a 9)

Cómo leer los veredictos

Cada veredicto es un juicio de este informe sobre un suceso concreto, con un plazo, 2031, y una condición, que se cumpla la premisa. No es una frecuencia ni sale de un modelo matemático: resume las razones de «Por qué estos veredictos», en cada capítulo, para que se pueda discrepar de ellas una a una.



Un tramo no es una cifra: «improbable» quiere decir entre el 5 y el 20 %, sin más. Los veredictos no se suman ni se multiplican, porque comparten causas: una crisis entre potencias traería a la vez ataques informáticos, sabotajes y escasez de chips. La gravedad es la del daño si ocurre, y no se promedia con la probabilidad. Un desenlace irreversible en el tramo más bajo puede merecer más trabajo que uno probable y recuperable, porque no deja corregir. Si la premisa no se cumple, casi todos bajan, y los que menos, los de los capítulos 4 a 6, que solo necesitan la mitad comprobable (capítulo 1). Y una cifra sin suceso ni plazo, como muchas de las que circulan sobre la extinción, no se compara con ninguna de estas (capítulo 9).

Referencias

Numeradas por orden de primera cita. Bajo cada entrada, el alcance del resultado que se utiliza. Consulta: septiembre de 2026.

- 1 Forecasting Research Institute. *Assessing Near-Term Accuracy in the Existential Risk Persuasion Tournament*. 2026. forecastingresearch.org ↗
Comprueba los pronósticos a corto plazo del torneo. La precisión a corto plazo no correlaciona con la estimación del riesgo a largo.
- 2 OpenAI. *GPT-6 Astra: A new generation of intelligence*. 2026-09-03. openai.com ↗
Resultados publicados por el desarrollador.
- 3 METR. *Frontier Risk Report*. 2026-05-19. metr.org ↗
Evaluación piloto de febrero-marzo, con acceso interno. Riesgo de despliegues no autorizados pequeños; no probabilidad calibrada de catástrofe.
- 4 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 ↗
Documento de trabajo; ola 1 (jun-ago 2025) con 330 expertos, 59 superpronosticadores y 1.360 del público. El 23 % es un pronóstico sobre lo que dirá el panel en 2030.
- 5 Dario Amodei. *The Adolescence of Technology: Confronting and Overcoming the Risks of Powerful AI*. 2026-01. darioamodei.com ↗
Ensayo de opinión; el autor advierte de que no pretende comunicar certeza ni probabilidad.
- 6 Dario Amodei. *We Must Pace the Frontier*. 2026-09. darioamodei.com ↗
Ensayo de opinión; propone frenar el ritmo de mejora de capacidades.
- 7 Softonic. *Anthropic's AI slowdown plan now has backing from Altman, Musk and Hassabis*. 2026-09. en.softonic.com ↗
Noticia; recoge los mensajes públicos de Altman, Musk y Hassabis sobre el ensayo del 12 de septiembre.
- 8 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 ↗

Transcripción oficial del pódcast; opiniones del director de SSI, parte interesada (su empresa apuesta por la investigación y no por el escalado). Texto guardado en 02_fuentes/ronda12.

- 9 Dario Amodei. *Machines of Loving Grace: How AI Could Transform the World for the Better*. 2024-10. darioamodei.com ↗
Ensayo del consejero delegado de Anthropic. Define «IA potente» y los factores que limitan a la inteligencia.
- 10 Help Net Security. *OpenAI just hit a milestone on the road to self-improving AI*. 2026-09-07. helpnetsecurity.com ↗
Cifras comunicadas por la propia empresa sobre su organización de investigación.
- 11 Sherwood News. *Google DeepMind's Hassabis: AGI is 3 to 4 years away*. 2026-05. sherwood.news ↗
Declaraciones de Demis Hassabis en Google I/O y a Axios; en junio de 2025 situaba la inteligencia general en 2030-2035.
- 12 AI Futures Project. *Clarifying how our AI timelines forecasts have changed since AI 2027*. 2025-11. blog.aifutures.org ↗
Revisión pública de los plazos de los autores de «AI 2027».
- 13 Medianama. *Sam Altman predicts superintelligence within years; Meta AI scientist disagrees*. 2026-02-21. medianama.com ↗
Crónica de la cumbre de Nueva Delhi (19 feb 2026) con citas literales de Altman («end of 2028») y de LeCun. Contrastado con Forbes India (19 feb 2026); no hay transcripción oficial localizada.
- 14 E. Erdil (Epoch AI). *The case for multi-decade AI timelines*. 2025-04-26. epoch.ai ↗
Artículo de opinión en el boletín de Epoch; extrapolación de ingresos y argumentos cualitativos. Mediana personal ~20 años hasta automatizar el trabajo a distancia.
- 15 Grace et al. (AI Impacts). *Thousands of AI Authors on the Future of AI*. 2024-01. arxiv.org ↗
Encuesta a 2.778 autores de IA; tasas de respuesta bajas y gran dispersión; probabilidades subjetivas.
- 16 The Decoder (sobre entrevista del Financial Times). *«You certainly don't tell a researcher like me what to do», says LeCun as he exits Meta*. 2026-01-03. the-decoder.com ↗
Resumen de una entrevista del FT (de pago, no abierta). «Dead end» es paráfrasis del FT.

- 17 TechCrunch (A. Heim). *Yann LeCun's AMI Labs raises \$1.03B to build world models*. 2026-03-09. techcrunch.com ↗
Noticia de prensa tecnológica seria; cifras de la ronda confirmadas también por Reuters según otros medios.
- 18 Anthropic. *Project Glasswing*. 2026-04-07. anthropic.com ↗
Anuncio del desarrollador. Las cifras de vulnerabilidades proceden de la propia empresa y de sus socios.
- 19 METR. *Task-Completion Time Horizons of Frontier AI Models (Time Horizon 1.1), datos publicados*. 2026-05-08. metr.org ↗
Horizonte: duración para un profesional humano de las tareas de software que el modelo resuelve con un 50 % o un 80 % de éxito. METR advierte de que por encima de 16 horas su batería deja de medir bien.
- 20 METR. *Impact of modeling assumptions on time horizon results*. 20-03-2026. metr.org ↗
Sensibilidad de la métrica cuando se satura la batería; horizonte de tareas no equivale a duración de autonomía fiable en cualquier trabajo.
- 21 OpenAI. *Research acceleration: The view inside OpenAI*. 2026-09-06. openai.com ↗
Cifras de la propia empresa sobre su organización de investigación. El consumo se valora a precios de venta de su API, no a coste interno, y la cifra de 600 dólares es la del investigador mediano.
- 22 Futurum Group. *AI Capex 2026: The \$690B Infrastructure Sprint*. 2026. futurumgroup.com ↗
Suma de las guías de inversión de Amazon, Alphabet, Microsoft, Meta y Oracle. Otras casas dan entre 630.000 y 800.000 millones según fecha y perímetro.
- 23 Epoch AI. *Trends in Artificial Intelligence*. 2026. epoch.ai ↗
Tendencias estimadas con intervalos de confianza amplios; se actualizan de forma continua.
- 24 CNBC. *OpenAI resets spending expectations, tells investors compute target is around \$600 billion by 2030*. 2026-02-20. cnbc.com ↗
Información de prensa sobre comunicaciones a inversores.
- 25 CNBC. *Anthropic tells investors annualized revenue run rate climbed to \$65 billion in July*. 2026-08-17. cnbc.com ↗
Ritmo anualizado comunicado a inversores; no son cuentas auditadas.
- 26 Cunningham, Althoff, Halperin, Jabarian, Koh, Ramani, Trammell, Whitfill y Wu. *The Economics of Recursive Self-Improvement*. 2026-07-13. elasticity.institute ↗

La capacidad se mide con el índice de capacidades de Epoch. La condición de aceleración autosostenida se cumple, en su calibración tentativa, si cada punto del índice eleva al menos un 15 % la productividad de la investigación; un cálculo «de servilleta» con datos de productividad de ingenieros da en torno a un 9 %. El umbral depende de otros parámetros inciertos.

- 27 METR. *The Economics of Recursive Self-Improvement*. 2026-07-22. metr.org ↗
Nota de dos investigadores de METR que presenta un artículo firmado por nueve economistas (elasticity.institute). Modelo simple y calibración «muy laxa», según sus autores.
- 28 Anthropic. *System Card: Claude Fable 5.1 & Claude Mythos 5.1*. 2026-09-01. www-cdn.anthropic.com ↗
Evaluación del propio desarrollador. Consultadas directamente las secciones de ciberseguridad (3.3 y 3.5), inyección de instrucciones (5.2) y alineamiento (6.2 y 6.3); el resto, a través de la reseña de Z. Mowshowitz.
- 29 Holmström y Milgrom. *Multitask Principal-Agent Analyses: Incentive Contracts, Asset Ownership, and Job Design*. 1991. [doi.org](https://doi.org/10.2307/2328686) ↗
Journal of Law, Economics, & Organization 7 (número especial). Premiar lo medible desvía la atención de lo que no se mide; modelo teórico.
- 30 Forecasting Research Institute. *LEAP Wave 1: Headliners*. 2025-11-10. leap.forecastingresearch.org ↗
Informe de la ola 1 con la pregunta literal y los resultados por grupo.
- 31 Epoch AI. *GATE: An Integrated Assessment Model for AI Automation*. 2025-03-21. arxiv.org ↗
Modelo macroeconómico con parámetros muy inciertos; los propios autores señalan supuestos irreales (sustitución perfecta sin robots). Autores individuales del arXiv no comprobados; citar como Epoch AI.
- 32 Anthropic. *Disrupting the first reported AI-orchestrated cyber espionage campaign*. 2025-11. anthropic.com ↗
Relato del propio proveedor sobre un abuso de su modelo; atribución con «alta confianza», no verificada por terceros.
- 33 AI Security Institute (Reino Unido). *How fast is autonomous AI cyber capability advancing?*. 2026-05-14. aisi.gov.uk ↗
Redes simuladas sin defensores activos; diez intentos por modelo; presupuestos de hasta 100 millones de tokens.
- 34 B. Joy. *Why the Future Doesn't Need Us*. 2000-04. wired.com ↗
Ensayo del cofundador de Sun Microsystems sobre nanotecnología, genética y robótica.

- 35 Cohen y Felson. *Social Change and Crime Rate Trends: A Routine Activity Approach*. 1979. [doi.org](https://doi.org/10.2307/2086845) ↗
American Sociological Review 44(4). Un delito exige que coincidan alguien con inclinación y capacidad, un objetivo adecuado y la falta de un guardián capaz. Pensada para delitos de contacto directo.
- 36 Anthropic. *Threat Intelligence Report, September 2026*. 2026-09. anthropic.com ↗
Casos seleccionados de abuso; visibilidad parcial, intención biológica ambigua en parte de los casos.
- 37 M. J. D. Vermeer, E. Lathrop y A. Moon (RAND). *On the Extinction Risk from Artificial Intelligence*. 2025. rand.org ↗
Informe RR-A3034-1; análisis exploratorio de escenarios, no probabilístico. Usado SOLO en sus vías nuclear y de geoingeniería; el capítulo biológico no se ha leído. Mes de publicación no comprobado (el artículo divulgativo es de mayo de 2025).
- 38 E. Erdil y M. Barnett (Epoch AI). *Most AI value will come from broad automation, not from R&D*. 2025-03-21. epoch.ai ↗
Ensayo con datos de productividad de EE. UU. (1988-2022); critica explícitamente el marco de Amodei.
- 39 A. Narayanan y S. Kapoor. *AI as Normal Technology*. 2025-04. knightcolumbia.org ↗
Ensayo; la tesis de la difusión lenta como amortiguador.
- 40 Wikipedia en inglés (comunidad de editores). *Wikipedia:Writing articles with large language models/RfC*. 2026-03-20. en.wikipedia.org ↗
Fuente primaria: cierre de la RfC (44 a 2); es una pauta (guideline); confirmado por TechCrunch (26 mar 2026), que da 40 a 2.
- 41 M. Schwartz y Z. Montague (The New York Times). *Artificial Intelligence floods court dockets with home-brewed lawsuits*. 2026-05-25. benningtonbanner.com ↗
Leído en reproducción sindicada; cita del juez Schiltz.
- 42 M. Spence. *Job Market Signaling*. 1973. [doi.org](https://doi.org/10.2307/2282624) ↗
Quarterly Journal of Economics 87(3). Una señal informa solo si cuesta más a quien no tiene la cualidad.
- 43 L. Torvalds (vía Slashdot). *Linux 7.1-rc4 announcement: AI reports make security list almost entirely unmanageable*. 2026-05-17. linux.slashdot.org ↗
Cita literal del correo a la LKML; confirmado por Tom's Hardware y TechRadar.
- 44 InfoWorld. *Internet Bug Bounty program hits pause on payouts*. 2026-04-03. infoworld.com ↗

Cita el anuncio de HackerOne; confirmado también por Dark Reading.

- 45 D. Stenberg. *curl summer of bliss; What the bliss taught us*. 2026-06-15. daniel.haxx.se ↗
Pausa de informes del 1 jul al 3 ago 2026; balance en <https://daniel.haxx.se/blog/2026/08/03/what-the-bliss-taught-us/>
- 46 AISI. *How far behind the frontier are leading open-weight models on cyber?*. 2026. aisi.gov.uk ↗
Brecha en sus pruebas: 4-7 meses frente a 6-10 en 2025. Entornos sin todos los obstáculos de una defensa real.
- 47 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 ↗
Autoevaluación del propio laboratorio antes de publicar gpt-oss; retos CTF y redes simuladas, más simples que sistemas reales. Usado solo en su parte de ciberseguridad.
- 48 Hugging Face. *State of Open Models: Summer 2026 Observations*. 2026. huggingface.co ↗
Descargas y modelos derivados en una sola plataforma.
- 49 Center for a New American Security. *Adversarial Distillation*. 2026. cnas.org ↗
Las cifras de cuentas e intercambios proceden de las empresas afectadas.
- 50 *International AI Safety Report 2026*. 2026-02. internationalaisafetyreport.org ↗
Síntesis internacional; capacidades, propensión y despliegue. Su evaluación corresponde a febrero, no a septiembre.
- 51 Cornish y Clarke. *Opportunities, Precipitators and Criminal Decisions: A Reply to Wortley's Critique of Situational Crime Prevention*. 2003. popcenter.asu.edu ↗
Crime Prevention Studies 16, pp. 41-96. Origen de la clasificación de la prevención situacional en cinco grupos y 25 técnicas.
- 52 Epoch AI. *LLM inference prices have fallen rapidly but unequally across tasks*. 2025-03. epoch.ai ↗
Caída del precio necesario para alcanzar un nivel de rendimiento dado: entre 9 y 900 veces al año según la tarea; mediana, 50.
- 53 TIME (entrevista a Y. LeCun). *Meta's AI Chief Yann LeCun on AGI, Open-Source, and AI Risk*. 2024-02-13. time.com ↗
Entrevista, fuente primaria de sus argumentos contra el riesgo existencial y a favor de abrir los pesos. Anterior a su salida de Meta.

- 54 D. Hendrycks, E. Schmidt y A. Wang. *Superintelligence Strategy: Expert Version*. 2025-03-07. arxiv.org ↗
Documento estratégico (no revisado por pares); define MAIM. Autores con intereses en el sector (Schmidt, Wang).
- 55 NCSC. *Impact of AI on cyber threat from now to 2027*. 2025. ncsc.gov.uk ↗
Evaluación prospectiva, no registro de incidentes; automatización y desigualdad de capacidades defensivas.
- 56 VulnCheck. *State of Exploitation 2026*. 2026. vulncheck.com ↗
Vulnerabilidades con evidencia pública de explotación. No atribuye la velocidad a la IA.
- 57 Chainalysis. *Crypto Ransomware: 2026 Crypto Crime Report*. 2026-02. chainalysis.com ↗
Pagos identificados en cadenas de bloques; la cifra de 2025 se revisará al alza.
- 58 INCIBE. *Balance de ciberseguridad 2025*. 2026. incibe.es ↗
Incidentes gestionados por INCIBE-CERT; no es un censo de todos los ataques en España.
- 59 J. Hirshleifer. *From weakest-link to best-shot: The voluntary provision of public goods*. 1983. doi.org ↗
Public Choice 41(3). Bienes comunes que dependen de la aportación más baja (eslabón más débil) o de la más alta (mejor tiro). No trata de seguridad informática.
- 60 H. R. Varian. *System Reliability and Free Riding*. 2004. doi.org ↗
En Camp y Lewis (eds.), *Economics of Information Security*, Kluwer. Aplica a la fiabilidad informática las tecnologías de Hirshleifer; modelo teórico.
- 61 Sophos. *The State of Ransomware 2026*. 2026-07. sophos.com ↗
Encuesta a 2.158 responsables de TI y ciberseguridad; datos declarados por las organizaciones.
- 62 CrowdStrike. *2026 Global Threat Report*. 2026-02. crowdstrike.com ↗
Intrusiones observadas por un proveedor de seguridad.
- 63 OWASP. *LLM Prompt Injection Prevention Cheat Sheet*. Consulta 2026-09-20. cheatsheetseries.owasp.org ↗
Contenido no confiable, permisos de agentes y controles externos al modelo.
- 64 CISA y socios. *PRC state-sponsored actors compromise and maintain persistent access to US critical infrastructure*. 2024-02-07. cisa.gov ↗
Preposicionamiento estatal en infraestructura; no atribución a IA.
- 65 Cybersecurity Dive. *What we know so far about the hacking campaign against US water systems*. 2026-07. cybersecuritydive.com ↗

Prensa especializada; atribución a grupos vinculados a Irán según las autoridades, no confirmada por terceros.

- 66 IAEA. *Computer Security of Instrumentation and Control Systems at Nuclear Facilities, NSS 33-T*. 2018. nucleus-apps.iaea.org ↗

Sistemas digitales y defensa en profundidad; no certifica instalaciones individuales.

- 67 NHS England. *Synnovis cyber incident*. 2025-11-10. england.nhs.uk ↗

Interrupción de proveedor de patología y aplazamientos sanitarios; no ataque causado por IA.

- 68 HIPAA Journal. *Patient Death Linked to Ransomware Attack on NHS Pathology Provider*. 2025-06. hipaajournal.com ↗

Una muerte con el retraso de una analítica como factor contribuyente, según el hospital afectado.

- 69 Microsoft. *Helping our customers through the CrowdStrike outage*. 2024-07-20. blogs.microsoft.com ↗

8,5 millones de equipos Windows caídos por una actualización defectuosa. No fue un ataque.

- 70 BleepingComputer. *Jaguar Land Rover cyberattack cost the company over \$220 million*. 2025-11. bleepingcomputer.com ↗

Cinco semanas de producción parada; el Cyber Monitoring Centre británico estima 1.900 millones de libras de impacto en la economía.

- 71 FBI, Internet Crime Complaint Center. *2025 IC3 Annual Report*. 2026-04. ic3.gov ↗

Solo denuncias presentadas en Estados Unidos; la pérdida real es mayor. Serie 2021-2024 tomada de los informes anuales anteriores.

- 72 UK Finance. *Annual Fraud Report 2026 (press release)*. 2026-06. ukfinance.org.uk ↗

Patronal bancaria británica; datos de sus miembros; informe completo en PDF en la web de UK Finance

- 73 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 ↗

Science 390(6777); N = 76.977; revisado por pares; detalles numéricos tomados del preprint arXiv 2507.13919.

- 74 Sky TG24; Il Fatto Quotidiano. *Truffa con nome (e voce) del ministro Crosetto: cosa sappiamo*. 2025-02-08. tg24.sky.it ↗

Prensa italiana seria; recuperación de fondos en Il Fatto Quotidiano (12 feb 2025).

- 75 C. Herley (Microsoft Research). *Why Do Nigerian Scammers Say They Are from Nigeria?*. 2012. microsoft.com ↗
Workshop on the Economics of Information Security. Al estafador lo limita el coste de atender a quien contesta y nunca va a pagar.
- 76 Europol. *IOCTA 2026 – The evolving threat landscape: how encryption, proxies and AI are expanding cybercrime*. 2026-04-28. europol.europa.eu ↗
Evaluación policial cualitativa; leída en resumen de prensa (The Cyber Express).
- 77 FTC. *Prepared Statement of the Federal Trade Commission before the Joint Economic Committee: The Rising Scam Economy*. 2026-03-25. ftc.gov ↗
Fuente oficial; solo pérdidas denunciadas (infradenuncia grande); no desglosa la IA.
- 78 CNN. *Someone using AI to impersonate Marco Rubio contacted at least five people including foreign ministers, cable says*. 2025-07-08. cnn.com ↗
Basado en un cable del Departamento de Estado del 3 jul 2025; también Washington Post, Bloomberg y NPR.
- 79 SRF. *Schwyz: Kriminelle erbeuten mit Künstlicher Intelligenz Millionen*. 2026-01-19. srf.ch ↗
Televisión pública suiza sobre comunicado de la Kantonspolizei Schwyz; importe exacto no publicado.
- 80 FinCEN (Departamento del Tesoro de EE. UU.). *FinCEN Alert on Fraud Schemes Involving Deepfake Media Targeting Financial Institutions (FIN-2024-Alert004)*. 2024-11-13. fincen.gov ↗
Alerta oficial; constata aumento de informes de operaciones sospechosas, sin cifras.
- 81 Entrust/Onfido (vía Infosecurity Magazine). *2025 Identity Fraud Report*. 2024-11-20. infosecurity-magazine.com ↗
Datos de un proveedor de verificación con interés comercial; millones de verificaciones en 195 países.
- 82 TechCrunch. *ID verification giant IDScan confirms data breach with more than 150 million driver's licenses stolen*. 2026-09-10. techcrunch.com ↗
Confirmado por la empresa; Krebs on Security y Time también; cifra exacta pendiente de investigación.
- 83 StartupHub.ai. *Geoffrey Hinton: 10 % chance AI kills humans*. 2026-09-11. startuphub.ai ↗
Medio menor; recoge una entrevista en BBC Politics sin fecha exacta. Coincide con Digital Today (Corea), 11 sep 2026. Fiabilidad baja-media: buscar el vídeo de la BBC antes de citar con comillas.

- 84 H. Lin, G. Czarnek, B. Lewis, J. P. White, A. J. Berinsky, T. Costello, G. Pennycook y D. G. Rand. *Persuading voters using human-artificial intelligence dialogues*. 2025-12-04. [nature.com](https://www.nature.com) ↗
Nature; experimentos preinscritos en tres elecciones; cifras por país según la nota de Cornell.
- 85 Cornell Chronicle. *AI chatbots can effectively sway voters – in either direction*. 2025-12. news.cornell.edu ↗
Nota de la universidad de Rand; resume Nature y Science con cifras.
- 86 L. Hölbling, S. Maier y S. Feuerriegel. *The Persuasive Power of LLMs: A Systematic Literature Review and Meta-Analysis*. 2025-09-08. doi.org ↗
Preprint (Research Square); solo 8 estudios; sirve de contrapeso.
- 87 Salvi et al.. *On the conversational persuasiveness of GPT-4*. 2025. doi.org ↗
Experimento de persuasión con 900 participantes; no prueba de efecto electoral agregado.
- 88 Salvi et al.. *Author Correction: On the conversational persuasiveness of GPT-4*. 2026-09-03. doi.org ↗
Corrección del contraste personalizado frente a no personalizado: no alcanza significación estadística; no invalida el contraste principal frente a humanos.
- 89 The Decoder. *Researchers used AI to manipulate Reddit users, scrapped study after backlash*. 2025-04. the-decoder.com ↗
Resultados del borrador no publicado (sin revisión); confirmado por Science News y Live Science.
- 90 K. Hackenburg, C. Wagner, L. Hewitt, B. M. Tappin, E. Saunders, H. R. Kirk, H. Margetts y C. Summerfield. *AI systems out-persuade expert humans*. 2026-06-15. arxiv.org ↗
Preprint sin revisión por pares; 6.923 personas; incluye prueba real de recaudación.
- 91 Common Sense Media. *Talk, Trust, and Trade-Offs: How and Why Teens Use AI Companions*. 2025-07. commonsensemedia.org ↗
Encuesta a 1.060 adolescentes de Estados Unidos (13-17 años).
- 92 OpenAI. *Strengthening ChatGPT's responses in sensitive conversations*. 2025-10-27. openai.com ↗
Estimaciones internas del proveedor sobre señales en conversaciones; difíciles de medir y no auditadas.
- 93 TechCrunch. *OpenAI says over a million people talk to ChatGPT about suicide weekly*. 2025-10-27. techcrunch.com ↗

- Reproduce las cifras de OpenAI (clave openai-sensibles); también Wired, 28 oct 2025.
- 94 M. y M. Raine contra OpenAI (Tribunal Superior de San Francisco, CGC-25-628528). *Complaint; OpenAI Answer (25 nov 2025)*. 2025-08-26. [courthousenews.com](https://www.courthousenews.com) ↗
Documentos judiciales: alegaciones de parte, sin sentencia; respuesta de OpenAI en <https://cdn.arstechnica.net/wp-content/uploads/2025/11/Raine-v-OpenAI-Answer-11-25-25.pdf>
 - 95 Social Media Victims Law Center y Tech Justice Law Project. *Lawsuits accuse ChatGPT of emotional manipulation, supercharging AI delusions, and acting as a «suicide coach»*. 2025-11-06. socialmediavictims.org ↗
Nota de los despachos demandantes (parte interesada); confirmado por el NYT.
 - 96 CBS News. *AI company, Google settle lawsuit over Florida teen's suicide linked to Character.AI chatbot*. 2026-01-07. [cbsnews.com](https://www.cbsnews.com) ↗
Confirmado por CNN, CNBC, Washington Post y Axios; términos no revelados.
 - 97 Legislatura de California. *SB 243 Companion chatbots*. 2025-10-13. leginfo.ca.gov ↗
Texto legal; en vigor desde el 1 ene 2026.
 - 98 MIT Media Lab. *How AI and Human Behaviors Shape Psychosocial Effects of Chatbot Use*. 2025. media.mit.edu ↗
Cuatro semanas, 981 participantes; no confundir asociación por uso voluntario con efecto causal de dosis.
 - 99 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 ↗
Documento de trabajo sin revisión por pares; >4,5 M de casos federales; el detector de texto de IA tiene error propio.
 - 100 C. Schmitz, L. Hammond y A. Chan. *Characterizing Agentic Flooding of Government Services*. 2026-08-17. arxiv.org ↗
Preprint (Hertie School, Cooperative AI Foundation, GovAI) para AIES 2026; 84 casos posibles en 11 jurisdicciones; muchos casos cualitativos.
 - 101 Aktuelle Sozialpolitik (con SWR e informe anual del LSG NRW). *Und täglich grüßt die «KI-Schleife»? KI bewegt die Sozialgerichte*. 2026-07-17. aktuelle-sozialpolitik.de ↗
Blog especializado que cita el informe anual 2025 del Landessozialgericht NRW y al LSG Renania-Palatinado.
 - 102 Cronkite News. *As more lawyers fall for AI hallucinations, ChatGPT says: Check my work*. 2025-10-28. cronkitenews.azpbs.org ↗

Periodismo universitario (Arizona State); cita directa de Charlotin y recuento de 486 casos.

- 103 D. Charlotin (HEC Paris). *AI Hallucination Cases Database*. 2026-09-21. damiencharlotin.com ↗
Base de datos primaria mantenida a mano; solo resoluciones que constatan la alucinación; es un suelo, no un censo.
- 104 A. Zahavi. *Mate selection—a selection for a handicap*. 1975. [doi.org](https://doi.org/10.1016/0022-5193(75)90047-0) ↗
Journal of Theoretical Biology 53(1). Principio del handicap: un rasgo costoso prueba la calidad. Formalizado por Grafen (1990); hoy discutido (Penn y Számadó, 2020).
- 105 Pangram Labs. *Pangram Predicts 21 % of ICLR Reviews are AI-Generated*. 2025-11-18. pangram.com ↗
Empresa que vende detectores (interés comercial); falsos positivos declarados 1/10.000; recogido por Nature News.
- 106 D. Stenberg. *The end of the curl bug-bounty*. 2026-01-26. daniel.haxx.se ↗
Blog del mantenedor principal; cifras propias del programa.
- 107 The Register. *Curl shuts bug bounty program to stop AI slop*. 2026-01-21. theregister.com ↗
Medio técnico serio; recuento de enero de 2026.
- 108 GPTZero. *GPTZero finds 100 new hallucinations in NeurIPS 2025 accepted papers*. 2026-01. gptzero.me ↗
Empresa que vende detectores; cada cita marcada se comprobó a mano según la empresa.
- 109 D. Stenberg. *The pressure*. 2026-05-26. daniel.haxx.se ↗
Blog del mantenedor; volumen de informes relativo a 2024 y 2025.
- 110 R. Brandom (TechCrunch). *AI agents are flooding public services with new requests*. 2026-09-10. techcrunch.com ↗
Periodismo; cifras del Housing Ombudsman y del CFPB; la atribución a la IA es de coincidencia temporal.
- 111 Financial Ombudsman Service. *Embracing AI's transformational impact on consumer complaints*. 2026. financial-ombudsman.org.uk ↗
Fuente oficial; muestra pequeña y no descrita; confirmado por Which? (15 may 2026).
- 112 RICS Land Journal. *What impact is AI having on planning systems?*. 2026. www3.rics.org ↗
Revista profesional de tasadores; testimonios de promotores, no datos oficiales.

-
- 113 The Decoder (con el Financial Times). *AI fraud forces world's largest accounting body to end online exams*. 2025-12-29. the-decoder.com ↗
Recoge la entrevista del FT a Helen Brand; confirmado por The Accountant.
- 114 The New York Times (vía eWeek). *Employers Are Buried in A.I.-Generated Résumés*. 2025-06-21. eweek.com ↗
Cifra de LinkedIn citada por el NYT; leída en eWeek, que la reproduce.
- 115 Greenhouse. *An AI Trust Crisis: 2025 AI in Hiring Report*. 2025-11-19. greenhouse.com ↗
Encuesta de una empresa de software de selección (4.136 personas, cuatro países); autodeclarado.
- 116 CNBC. *How Google is responding to AI cheating in coder interviews*. 2025-03-09. cnbc.com ↗
Periodismo con audio de reunión interna de Google; Pichai lo confirmó en jun 2025 (Lex Fridman).
- 117 arXiv. *Attention Authors: Updated Practice for Review Articles and Position Papers in arXiv CS Category*. 2025-10-31. blog.arxiv.org ↗
Anuncio oficial de arXiv.
- 118 TechCrunch. *Research repository ArXiv will ban authors for a year if they let AI do all the work*. 2026-05-16. techcrunch.com ↗
Cita el anuncio de T. Dietterich; confirmado por Nature (d41586-026-01595-5).
- 119 Naciones Unidas. *General Assembly Adopts More Than 60 Resolutions, Decisions of Its First Committee (GA/12736)*. 2025-12-01. press.un.org ↗
Resolución 80/57 sobre sistemas de armas autónomas letales, aprobada en el pleno por 164 votos a favor, 6 en contra (Bielorrusia, Burundi, Corea del Norte, Israel, Rusia y Estados Unidos) y 7 abstenciones. No es vinculante.
- 120 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 ↗
Estimaciones de analistas (SemiAnalysis, Goldman Sachs); los rendimientos de fabricación no son públicos.
- 121 Oklahoma Watch. *Have 70 % of casualties in the Russia-Ukraine war been caused by drones?*. 2026-06-09. oklahomawatch.org ↗
Verificación periodística de la cifra; procede de estimaciones ucranianas y occidentales no comprobables de forma independiente.
- 122 Bueno de Mesquita, Smith, Siverson y Morrow. *The Logic of Political Survival*. 2003. doi.org ↗

- MIT Press. Teoría del selectorado. Su prueba estadística está discutida (Clarke y Stone, 2008).
- 123 Beblawi y Luciani (eds.). *The Rentier State*. 1987. [iai.it](#) ↗
Croom Helm. Teoría del Estado rentista. La relación empírica entre petróleo y autoritarismo está discutida (Ross 2001; Haber y Menaldo 2011).
- 124 Kulveit et al.. *Gradual Disempowerment: Systemic Existential Risks from Incremental AI Development*. 2025-01. [arxiv.org](#) ↗
Argumento teórico, no estimación empírica.
- 125 Drago y Laine. *The Intelligence Curse*. 2025-04. [intelligence-curse.ai](#) ↗
Serie de ensayos sin revisión por pares. Aplica a la IA la maldición de los recursos: quien vive de la IA pierde el incentivo a invertir en las personas.
- 126 Human Rights Watch. *UN Talks on Killer Robots End with Calls for Negotiations Growing*. 2026-09-07. [hrw.org](#) ↗
Organización que hace campaña por un tratado.
- 127 The Next Web, con datos del Stanford AI Index 2026. *Stanford AI Index 2026: China narrows US lead to 2.7 % while spending 23x less on AI investment*. 2026. [thenextweb.com](#) ↗
La brecha se mide en una clasificación de preferencias de usuarios; la inversión privada china no incluye fondos estatales.
- 128 Boston Consulting Group y Semiconductor Industry Association. *Strengthening the Global Semiconductor Supply Chain in an Uncertain Era*. 2021-04. [semiconductors.org](#) ↗
Dato de 2019-2021: el 92 % de la capacidad de lógica por debajo de 10 nm estaba en Taiwán. Las fábricas de Arizona y Japón lo han reducido algo desde entonces.
- 129 Office of the Director of National Intelligence. *Annual Threat Assessment of the U.S. Intelligence Community 2026*. 2026-03. [intelligence.senate.gov](#) ↗
Evaluación no clasificada; criticada por algunos analistas por restar gravedad al riesgo.
- 130 CSIS. *Lights Out? Wargaming a Chinese Blockade of Taiwan*. 2025-07-31. [csis.org](#) ↗
Simulación de escenarios; no probabilidad de guerra ni predicción puntual.
- 131 Bloomberg Economics. *Xi, Biden and the \$10 Trillion Cost of War Over Taiwan*. 2024-01-09. [bloomberg.com](#) ↗
Estimación de modelo: guerra $\approx$ 10 % del PIB mundial; bloqueo $\approx$ 5 %.
- 132 Tom's Hardware. *Chinese state media counters Anthropic's call to put brakes on AI development*. 2026-09. [tomshardware.com](#) ↗

Reacción de medios estatales y del Ministerio de Exteriores.

- 133 D. Abecassis (MIRI). *Refining MAIM: Identifying Changes Required to Meet Conditions for Deterrence*. 2025-04-11. intelligence.org ↗
Crítica desde MIRI (favorable a frenar, pero escéptica de la disuasión por sabotaje).
- 134 J. R. Arnold (AI Frontiers). *Superintelligence Deterrence Has an Observability Problem*. 2025-08-14. ai-frontiers.org ↗
Ensayo en AI Frontiers (revista vinculada, según creo, al Center for AI Safety de Hendrycks; vínculo no comprobado).
- 135 SIPRI. *The Impact of Military Artificial Intelligence on Nuclear Escalation Risk*. 2025-06. sipri.org ↗
Mecanismos de escalada, percepción y decisión; no frecuencia estimada de guerra nuclear.
- 136 Anthropic. *Agentic Misalignment in Summer 2026*. 2026. alignment.anthropic.com ↗
Búsqueda deliberada de conductas problemáticas en escenarios experimentales; no tasa de incidentes en uso ordinario.
- 137 METR. *Investigation of the OpenAI Hugging Face incident*. 2026-08-26. metr.org ↗
Investigación acotada de coordinación y manipulación de evaluaciones; no auditoría completa de infraestructura ni prueba de toma de control.
- 138 E. Yudkowsky, N. Soares y D. Sabien. *If Anyone Builds It, Everyone Dies: One Year Closer*. 2026-09-16. lesswrong.com ↗
Entrada de aniversario de los autores (MIRI); fuente primaria de las seis tesis del libro. Su lectura de los incidentes de 2026 es interpretación de parte.
- 139 J. Carlsmith. *Is Power-Seeking AI an Existential Risk?*. 2022-06-16. arxiv.org ↗
Informe de Open Philanthropy (borrador público abr 2021; v2 ago 2024). Probabilidades subjetivas e inestables según el propio autor: ~5 % en 2021, >10 % desde mayo de 2022. El autor trabaja en Anthropic desde finales de 2025.
- 140 C. A. E. Goodhart. *Problems of Monetary Management: The U.K. Experience*. 1975. rba.gov.au ↗
Ponencia de 1975, publicada en Papers in Monetary Economics I (Banco de la Reserva de Australia, 1976). Origen de la ley de Goodhart, formulada para la política monetaria.
- 141 M. Strathern. *'Improving ratings': audit in the British University system*. 1997. cambridge.org ↗

European Review 5(3), p. 308: formulación popular de la ley de Goodhart, tomada de Keith Hoskin.

- 142 Anthropic. *From shortcuts to sabotage: natural emergent misalignment from reward hacking*. 2025-11-21. anthropic.com ↗

Experimento de entrenamiento; generalización de atajos hacia otras conductas indeseadas.

- 143 Anthropic. *Agentic Misalignment: How LLMs could be insider threats*. 2025-06. anthropic.com ↗

Escenarios artificiales contruidos para forzar el dilema; no es una tasa en uso real.

- 144 OpenAI y Apollo Research. *Detecting and reducing scheming in AI models*. 2025-09. openai.com ↗

La mejora medida puede deberse en parte a que el modelo reconoce que está siendo evaluado.

- 145 W. R. Ashby. *An Introduction to Cybernetics*. 1956. pespmc1.vub.ac.be ↗

Chapman & Hall. Ley de la variedad necesaria (cap. 11): solo la variedad del regulador puede absorber la de las perturbaciones.

- 146 Scientific American (P. Hall). *AI researcher Jacob Coxon quit, fearing extinction. Security experts see a familiar fight*. 2026-09-11. scientificamerican.com ↗

Prensa científica seria; recoge la cifra de Hubinger y el contrapunto de expertos en seguridad (Trail of Bits, HackerOne, Kapoor).

- 147 W. MacAskill. *A short review of «If Anyone Builds It, Everyone Dies»*. 2025-09-18. willmacaskill.substack.com ↗

Reseña de un filósofo que se toma en serio el riesgo; crítica desde dentro del campo.

- 148 LawZero. *Yoshua Bengio launches LawZero: a new nonprofit advancing safe-by-design AI*. 2025-06-03. lawzero.org ↗

Nota de prensa de la propia organización; la cifra de 30 M\$ viene de TechCrunch y TIME, no de la nota.

- 149 BetaKit. *Yoshua Bengio's LawZero receives \$300 million backing from Canada and Germany*. 2026-09-16. betakit.com ↗

Prensa tecnológica canadiense; «hasta» 300 M de dólares canadienses según The Globe and Mail. Reparto entre los dos países no desglosado.

- 150 J. Carlsmith. *Superforecasting the premises in «Is power-seeking AI an existential risk?»*. 2023-10-18. joecarlsmith.com ↗

Resume las previsiones de 21 superpronosticadores de Good Judgment (agre-

- gado de mayo de 2023): 1 % frente a su 5 %. Encuesta financiada por Open Philanthropy.
- 151 Forecasting Research Institute. *Forecasting Existential Risks: Evidence from a Long-Run Forecasting Tournament*. 2023. forecastingresearch.org ↗
Torneo de 2022 con 169 participantes. Probabilidades subjetivas, no frecuencias.
 - 152 G. Kestin y N. Soares. *Views on AI Existential Risk Before and After a Public Event at Harvard University*. 2026-03-29. arxiv.org ↗
Preprint sin revisión por pares; coautor = conferenciante y coautor del libro promocionado (conflicto de intereses). n = 89 emparejados; subgrupos de 15 y 10. La media sube 6 puntos; la mediana, un tramo.
 - 153 TechCrunch (R. Bellan). «*Gambling with our lives*»: Anthropic researcher quits, warns against self-improving AI. 2026-09-09. techcrunch.com ↗
Noticia con citas literales del hilo de Coxon y de la respuesta de Hubinger en X. Opiniones personales, no posición de la empresa.
 - 154 Bloomberg (M. Barkley). *AI Pioneer Andrew Ng Calls Extinction Fears «Science Fiction»*. 2026-09-17. bloomberg.com ↗
De pago; citas contrastadas en The Next Web y Startup Fortune (17 sep 2026). «much more science fiction than science» es literal en todos; «wave of PR» viene del resumen de Bloomberg.
 - 155 C. Collier (Asterisk). *More Was Possible: A Review of If Anyone Builds It, Everyone Dies*. 2025-09. asteriskmag.com ↗
Reseña larga en revista del entorno de la seguridad de la IA; crítica desde dentro del campo.
 - 156 The Guardian. «*Godfather of AI*» raises odds of the technology wiping out humanity over next 30 years. 2024-12-27. theguardian.com ↗
Crónica de su entrevista en BBC Radio 4 (Today). No abierto directamente (bloqueado); citas contrastadas en Slashdot y Forbes. Autoría del artículo no comprobada.
 - 157 The Next Web (A. M. Constantin). *Andrew Ng calls the AI extinction warnings science fiction*. 2026-09-17. thenextweb.com ↗
Crónica de la entrevista de Bloomberg; atribuye a Ng la frase «pull up the ladder». La cifra de 2023 que da no está comprobada.
 - 158 AISI. *Frontier AI Trends Report*. 2025. aisi.gov.uk ↗
Evaluaciones de sistemas hasta octubre de 2025. Competencia en pruebas biológicas no equivale a daño real.
 - 159 RAND. *Can LLM Agents Select and Engage Biological Tools?*. 2026-06-25. rand.org ↗

Siete agentes; interacción inicial con herramientas, no producción de un agente dañino ni ataque completo.

- 160 RAND. *Building a Defense-in-Depth Biosecurity Strategy for the AI Era*. 2026-08-18. [rand.org](#) ↗
Propuesta de capas de prevención, detección y respuesta para distintos tipos de actores.
- 161 RAND. *The Operational Risks of AI in Large-Scale Biological Attacks*. 2024. [rand.org](#) ↗
Pruebas de planificación con modelos de 2023; no aumento estadísticamente significativo en esa prueba, no garantía actual.
- 162 WHO. *Laboratory biosecurity guidance*. 2024-06-21. [who.int](#) ↗
Bioseguridad institucional y protección de materiales, información y procesos.
- 163 WHO. *Global guidance framework for the responsible use of the life sciences*. 2022-09-13. [who.int](#) ↗
Doble uso, responsabilidad durante el ciclo de investigación y salud humana, animal y vegetal.
- 164 Meyer et al. (Microsoft). *How effective is multifactor authentication at deterring cyberattacks?*. 2023-05. [arxiv.org](#) ↗
Estudio observacional sobre cuentas de Azure Active Directory.
- 165 Lee et al.. *The Impact of Generative AI on Critical Thinking*. 2025. [microsoft.com](#) ↗
Encuesta a 319 trabajadores; esfuerzo cognitivo autodeclarado, no medición de pérdida de inteligencia.
- 166 L. Bainbridge. *Ironies of Automation*. 1983. [doi.org](#) ↗
Automatica, 19(6). Cuanto más se automatiza, más depende el sistema de un operador muy cualificado que ya no practica.
- 167 Wellerstein, A.. *Counting the dead at Hiroshima and Nagasaki*. 2020-08-04. [thebulletin.org](#) ↗
Bulletin of the Atomic Scientists.
- 168 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](#) ↗
Nota de prensa del informe anual de mortalidad en la niñez.
- 169 Fundación Gates. *Goalkeepers 2025: child deaths projected to rise for the first time this century*. 2025-12-04. [gatesfoundation.org](#) ↗
Proyección del IHME; niveles no comparables con los de la ONU.