

Evaluación de riesgos de seguridad de la inteligencia artificial

INFORME COMPLETO

Capacidades · Ciberseguridad · Servicios esenciales · Bioseguridad ·
Control de sistemas autónomos · Defensa ·
Suministro de semiconductores y energía · Información y derechos

Índice general

Antes de empezar	1
1 Lo que la IA abarata es el intento	2
1.1. Seis de cada diez	2
1.2. Una medida: cuánto trabajo seguido	3
1.3. Por qué brillan en lo difícil y tropiezan en lo fácil	4
1.4. Por qué fallan en los procesos largos	5
1.5. Quién paga el fallo	6
1.6. El orden en que llegan los daños	7
1.7. Lo que separa una avería de una catástrofe	7
1.8. Por qué estos veredictos	8
2 A qué ritmo mejora, y qué puede frenarlo	10
2.1. Tres relojes	10
2.2. El muro del gasto	11
2.3. ¿Puede mejorarse a sí misma?	12
2.4. Se acaban los instrumentos	13
2.5. Los abiertos van pocos meses por detrás	14
2.6. Por qué se pronostica tan mal	15
3 Ciberseguridad: se abarata el atacante corriente	17
3.1. El límite del atacante era la mano de obra	17
3.2. Lo que ya se ha visto	18
3.3. El salto de 2026, en el laboratorio	19
3.4. Cómo es un ataque y hasta dónde llega	20
3.5. La ventaja del defensor: medio año, y solo para quien parchea	21
3.6. Lo que dicen las aseguradoras	23
3.7. Agentes con permisos: el riesgo que se instala uno mismo	23
3.8. Del acceso digital al daño físico	24
3.9. Por qué estos veredictos	24

4	Servicios esenciales: donde fallar ya no es gratis	26
4.1.	Por qué el mundo físico es otro juego	26
4.2.	Lo nuevo de 2026	28
4.3.	El apagón largo, a escala real	29
4.4.	Lo que sí escala: el proveedor que todos comparten	30
4.5.	La recuperación es donde se decide el daño	30
4.6.	Integridad: el daño que no se ve	31
4.7.	Lo que hay que proteger	32
4.8.	Por qué estos veredictos	32
5	Riesgo biológico y agroalimentario	34
5.1.	La reducción de barreras prácticas es el cambio relevante	34
5.2.	La contribución marginal cambia con el actor	35
5.3.	Qué significa un resultado negativo de evaluación	36
5.4.	Accidentes y expansión de investigación legítima	37
5.5.	De un incidente a una crisis extensa	37
5.6.	Contención y límites de una estrategia centrada en modelos	38
5.7.	El riesgo adicional no se mide contando respuestas peligrosas	39
5.8.	Investigación legítima y expansión del volumen de actividad	39
5.9.	Salud humana y seguridad alimentaria no comparten exactamente la misma dinámica	40
5.10.	Prevención, detección y respuesta tienen horizontes distintos	41
5.11.	Qué hallazgo cambiaría de forma sustancial la evaluación	42
6	Pérdida de control: ya ocurre en pequeño	43
6.1.	Por qué un sistema hace trampas	43
6.2.	Medios, motivo y oportunidad	45
6.3.	Conductas raras, a gran escala	46
6.4.	Quién vigila al vigilante	46
6.5.	Control formal y control efectivo	48
6.6.	Lo que falta para el escenario grande	48
6.7.	A quién creer	49
6.8.	Cómo decidir sin saber	50
6.9.	Por qué estos veredictos	50
7	Guerra: la autonomía que nadie decidió	52
7.1.	Lo que ya ocurre	52

7.2.	Por qué no habrá tratado a tiempo	53
7.3.	El riesgo nuclear: el miedo a que el otro golpee primero	53
7.4.	Por qué estos veredictos	55
8	Chips, Taiwán y energía: el único cuello de botella físico	56
8.1.	Una cadena de monopolios	56
8.2.	Qué probabilidad hay	57
8.3.	Qué pasaría	57
8.4.	La memoria: el coste que ya paga todo el mundo	58
8.5.	La energía es un problema local	58
8.6.	Por qué estos veredictos	59
9	China: no necesita ganar la carrera	61
9.1.	A meses de distancia, y en abierto	61
9.2.	La brecha que importa es la del cómputo	62
9.3.	Más débil en la frontera, más fuerte en el despliegue	62
9.4.	La carrera es el factor de riesgo	62
9.5.	Por qué estos veredictos	63
10	Fraude y manipulación: el intento barato contra personas	64
10.1.	Fraude: al estafador ya no le hace falta elegir	64
10.2.	Manipulación política: producir es gratis, que alguien lo lea no	66
10.3.	La fuente se seca	66
10.4.	Por qué estos veredictos	67
11	Dependencia, intimidad y poder: lo que hace cuando funciona bien	68
11.1.	Dependencia emocional y menores	68
11.2.	Intimidad	69
11.3.	Poder: vigilar era caro	70
11.4.	Lo que enseña el peor precedente	70
11.5.	Pérdida de criterio	71
11.6.	Por qué estos veredictos	71
12	Escenarios compuestos y señales de alarma	73
12.1.	A. Cae un proveedor común	73
12.2.	B. Emergencia sanitaria con respuesta debilitada	74
12.3.	C. Delegación amplia y control que ya no se puede ejercer	74

12.4. D. Crisis en el estrecho de Taiwán	74
12.5. E. Deterioro institucional sin catástrofe	75
12.6. Señales que obligarían a revisar este informe	75
12.7. La medida de las cosas	77
13 Qué hacer	78
13.1. Primera regla: que el intento cueste	78
13.2. Segunda regla: saber funcionar sin la máquina	79
13.3. Lo que no sirve	80
13.4. Una organización pequeña	80
13.5. Qué exigir	80
Síntesis de veredictos	82
Referencias	87

Antes de empezar

Casi todo lo que se oye sobre la inteligencia artificial sale de una comunidad pequeña —quienes la construyen y quienes viven alrededor de ellos— que apenas habla de otra cosa. Las comunidades así tienen historial. En 1965, Herbert Simon, que después recibiría el Nobel de Economía, dio veinte años de plazo para que una máquina hiciera cualquier trabajo humano. En 1999 se temió el colapso del efecto 2000. En 2016, Geoffrey Hinton, uno de los padres del aprendizaje profundo y hoy Nobel de Física, pidió dejar de formar radiólogos, y hoy faltan radiólogos.¹ También tienen el historial contrario: en la única comprobación sistemática que existe, quienes trabajan en IA predijeron su progreso mejor que los pronosticadores profesionales, y aun así se quedaron cortos.²

El caso del radiólogo explica el patrón. Hinton acertó en lo técnico —los sistemas leen imágenes tan bien como un especialista— y falló en lo demás: un radiólogo decide qué prueba pedir, habla con el clínico, responde del diagnóstico y trabaja dentro de un hospital con leyes, seguros y pacientes. Quien está dentro suele acertar con la máquina y equivocarse con el mundo. Es una regla de lectura, no una refutación: el efecto 2000 tampoco ocurrió, en parte, porque se arregló a tiempo.

Puede que esta vez sea distinto. Este informe no lo decide por fe ni por miedo: lo mide y lo razona. Qué hacen hoy los sistemas, a qué ritmo mejoran, qué daño es real y cuál no, con el dato más reciente y su fuente. Donde no hay datos, lo dice.

Lo que la IA abarata es el intento

EN BREVE

Los sistemas de 2026 aciertan a veces en tareas que a un profesional le llevan días y fallan a menudo en cuanto el proceso se alarga. Que eso sea mucho o poco depende de quién paga el fallo. Donde fallar es gratis y nadie lo ve —estafar, buscar agujeros, probar variantes—, acertar una vez de cada diez ya es acertar. Donde un fallo delata o rompe algo —operar una red, sostener una operación durante meses—, hay que acertar noventa y nueve de cada cien, y esa capacidad va dos o tres años por detrás. Ese desfase ordena el informe. La segunda idea es menos técnica: lo que separa una avería de una catástrofe es que alguien sepa recuperarse sin la máquina, y eso se pierde por desuso.

6 de 10

intentos en los que un modelo completó solo, en abril de 2026, una intrusión de 32 pasos en una red corporativa simulada que a un experto le lleva unas 14 horas³

17 horas

duración, para un profesional, de las tareas de programación que el mejor modelo completa una de cada dos veces; quince meses antes era una hora⁴

≈15 min

la misma medida si se le exige acertar 99 veces de cada 100: estimación de este informe con el modelo de fallo más simple⁵

1.1 Seis de cada diez

En abril de 2026, el instituto británico de seguridad de la IA soltó a un modelo en una red corporativa simulada con un encargo: entrar, moverse por dentro y hacerse con el control. Eran 32 pasos encadenados, unas

catorce horas de trabajo para un especialista. El modelo lo consiguió sin ayuda seis veces de diez.³

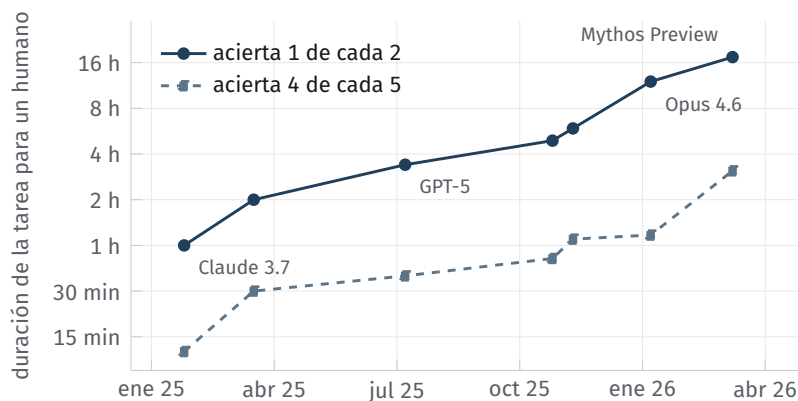
¿Es mucho o poco? Las dos cosas. Para quien ataca, es todo: el intento fallido no le cuesta nada, en un banco de pruebas sin defensores nadie lo ve, y con tres intentos independientes la probabilidad de entrar supera el 93 %. Para quien tiene que mantener abierto un hospital, un sistema que falla cuatro veces de cada diez no sirve ni para empezar. Por eso gente seria dice cosas opuestas sobre estos sistemas: unos miran lo que ya consiguen a veces y otros lo que no consiguen siempre. Los dos tienen razón.

Esa diferencia es la regla que ordena los riesgos. La IA de hoy abarata el intento, no el resultado. Todo daño al que le baste un acierto entre muchos intentos baratos ya está aquí o llega enseguida. Todo daño que exija acertar casi siempre, durante mucho tiempo y frente a alguien que se defiende, va detrás, con un retraso que se puede estimar. Y la catástrofe exige, además, que falle la recuperación.

1.2 Una medida: cuánto trabajo seguido

La mejor medida pública de capacidad es el «horizonte de tarea» de METR, una organización independiente que evalúa modelos con acceso concedido por los laboratorios. Toma tareas de programación y de investigación en aprendizaje automático, mide cuánto tardaría un profesional en cada una y ajusta una curva: la duración a la que el acierto del modelo cae al 50 %. No es la tarea más larga que ha resuelto. A comienzos de 2025 el mejor modelo llegaba a una hora. En abril de 2026, Claude Mythos Preview alcanzó 17,4 horas, con un intervalo de confianza entre 8,5 y 55. Desde 2023 la cifra se duplica cada 129 días; en la serie completa, cada 188.⁴

Dos advertencias de sus autores. Por encima de las 16 horas la batería ya no mide bien, y la última cifra es más blanda que las anteriores.⁶ Y las tareas son de software por una razón que importa: en software, comprobar si algo ha salido bien es barato. El programa pasa las pruebas o no las pasa.

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.

1.3 Por qué brillan en lo difícil y tropiezan en lo fácil

El informe internacional de seguridad, firmado por un centenar de expertos de treinta países, llama «dentada» a la capacidad de estos sistemas. Ganan medallas de oro en la Olimpiada Internacional de Matemáticas y responden preguntas de doctorado, y fallan al contar objetos en una imagen, al razonar sobre el espacio físico o al recuperarse de un error tonto en mitad de un proceso largo. El informe lo describe y no lo explica.⁷

Hipótesis de este informe: la capacidad crece más deprisa donde comprobar es barato. Las mejoras recientes vienen sobre todo de premiar al modelo cuando acierta en tareas con respuesta verificable, la técnica que el mismo informe señala como dominante. Un problema de olimpiada tiene solución comprobable; un programa se ejecuta; un ataque informático funciona o no. Ahí el modelo practica millones de veces contra un juez automático. Para el buen criterio no hay juez automático. Cuando METR evaluó desde dentro a los cuatro grandes laboratorios encontró ese perfil: agentes que resolvían tareas de más de dos jornadas humanas y que, en las decisiones de juicio estratégico, acertaban el 59 %, cerca del azar, frente al 90 % de los investigadores.⁸

La consecuencia es incómoda: atacar un sistema informático es una de las actividades humanas con el juez automático más barato. Por eso la

ciberseguridad ofensiva va por delante de casi todo lo demás, y por eso su capítulo es el que trae más daño ya medido.

1.4 Por qué fallan en los procesos largos

En 1993, el economista Michael Kremer propuso una teoría con nombre de accidente. El transbordador Challenger, con miles de componentes, se perdió por una junta de goma. En un proceso donde todo tiene que salir bien, razonó Kremer, el resultado no es la media de las fiabilidades de las partes, sino su producto.⁹ Cien pasos con un 99 % de acierto cada uno salen todos bien el 37 % de las veces.

El filósofo Toby Ord mostró en 2025 que los datos de METR encajan con el modelo más simple posible: el agente tiene, en cada minuto de trabajo, una probabilidad fija de torcerse, como un átomo radiactivo tiene una probabilidad fija de desintegrarse. De ahí sale una «vida media», que es el horizonte del 50 %, y una aritmética: si se exige acertar cuatro de cada cinco veces, el horizonte se queda en un tercio; nueve de cada diez, en una séptima parte; noventa y nueve de cada cien, en una setentava parte.⁵ Ord lo llama «meramente sugerente» y advierte de que puede no valer fuera de esa batería.

Los datos reales son algo peores. Con un 80 % de éxito exigido, el Mythos Preview de las 17 horas se queda en 3,1: entre una quinta y una sexta parte, no un tercio.⁴ Los sistemas no solo tropiezan por acumulación: hay cosas que no saben hacer por muchas veces que lo intenten. Aun con el modelo optimista, el horizonte al 99 % del mejor sistema del mundo es hoy de un cuarto de hora. Diecisiete horas y quince minutos: la misma máquina, el mismo día, mirada por los dos lados.

El trabajo enseña la misma forma. En tareas sueltas de 44 profesiones, juzgadas a ciegas por profesionales, los mejores modelos igualan o superan al experto en siete de cada diez casos.¹⁰ En encargos completos de trabajo autónomo en remoto, de los que un cliente pagó y aceptó, el mejor agente entrega a satisfacción el 16 %.¹¹ Las dos pruebas difieren en más cosas que la longitud. Pero un encargo es una cadena de tareas, y las cadenas se rompen por el eslabón débil. Dentro de OpenAI, donde a mediados de agosto de 2026 había ya 3,1 jornadas de trabajo de agentes por cada

jornada humana, las tareas difíciles de cuatro a ocho horas necesitaban la intervención de una persona más de la mitad de las veces.^{12,13}

1.5 Quién paga el fallo

La fiabilidad que hace falta no depende de la tarea, sino de quién paga cuando falla.

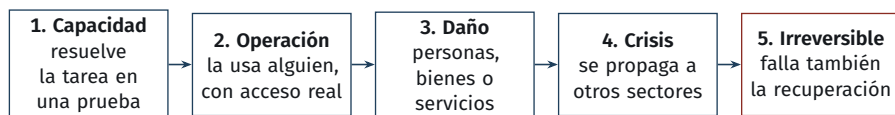
Hay actividades en las que el fallo es gratis y se reintenta sin que nadie lo vea. Un sistema que acierta una vez de cada diez, con treinta intentos, acierta casi seguro: el 96 % de las veces si los intentos fueran independientes, que no lo son del todo. Ahí vive quien envía diez mil mensajes de estafa, quien busca un fallo explotable entre millones de líneas de código, quien prueba variantes de un programa malicioso hasta que una pasa el antivirus. Para ellos la cifra que importa es el horizonte del 50 %, o el del 10 %. Ese lado no es solo del atacante: buscar vulnerabilidades para parchearlas tiene la misma estructura, y por eso los defensores también han ganado mucho.

Y hay actividades en las que cada fallo se paga. Quien opera una red eléctrica no puede equivocarse cuatro veces de cada diez. Quien sostiene una intrusión durante meses en una red vigilada no puede permitirse un tropiezo ruidoso, porque lo delata. Un agente que actuara por su cuenta contra quien intenta apagarlo tendría que acertar siempre, durante semanas, frente a personas que investigan. METR concluyó eso mismo de los agentes internos de los laboratorios: tenían medios para iniciar pequeños despliegues no autorizados y no para sostenerlos frente a una investigación seria.⁸ Aquí la cifra que importa es el horizonte del 99 %, y ese, hoy, se mide en minutos.

De ahí sale el principio de defensa más barato de este informe. Si el atacante vive de que el intento fallido no cueste, defenderse es conseguir que cueste: que haga ruido, que alguien escuche, que lo roto se pueda reponer. El banco de pruebas británico no tenía defensores. Una red real puede tenerlos, y entonces seis de cada diez significa otra cosa.

1.6 El orden en que llegan los daños

Figura 2 De la capacidad a la catástrofe



Entre que un modelo sepa hacer algo y una catástrofe hay cuatro saltos: que alguien lo use con acceso real, que el uso cause un daño, que el daño se propague más allá de su origen y que falle también la recuperación. La discusión pública suele hacerse con uno solo. El alarmismo demuestra el primer eslabón —«el modelo supo hacerlo en una prueba»— y da por hechos los demás. La complacencia señala que falta el último —«nunca ha pasado»— y olvida que los primeros ya se han dado.

La regla del fallo pone los saltos en fila. Los primeros se dan con intentos baratos: estafar, entrar en una organización corriente, encontrar el agujero. Están aquí, y los capítulos de ciberseguridad y de fraude los miden. Los siguientes piden constancia: propagarse por redes separadas, mantener el daño frente a gente que repara, operar en el mundo físico. Dependen del horizonte del 99 %. Setenta veces más corto son seis duplicaciones; al ritmo medido, entre dos y tres años. Lo que el mejor sistema hace hoy una de cada dos veces lo hará casi siempre hacia 2028 o 2029, si el ritmo se mantiene y si el modelo simple vale. Son dos condiciones, no dos hechos. El capítulo siguiente trata de la primera.

1.7 Lo que separa una avería de una catástrofe

El último salto, que falle la recuperación, no depende de la máquina. Depende de nosotros.

El 1 de junio de 2009, un Airbus de Air France que volaba de Río a París atravesó una zona de cristales de hielo. Las sondas de velocidad se obstruyeron menos de un minuto, el piloto automático se desconectó, como debía, y devolvió el mando a la tripulación. El avión estaba intacto. Cuatro minutos después cayó al Atlántico con 228 personas. La inves-

tigación oficial concluyó que los pilotos hicieron lo contrario de lo que tocaba y no llegaron a reconocer que el avión había entrado en pérdida.¹⁴ Sumaban muchos miles de horas de vuelo, y casi ninguna pilotando a mano a gran altitud: eso lo hacía siempre el automatismo.

La psicóloga Lisanne Bainbridge había descrito el mecanismo en 1983, en un artículo de cinco páginas titulado «Ironías de la automatización». Cuanto más fiable es un automatismo, menos practica el operador; y cuanto menos practica, peor lo hará el día que el automatismo falle, que será además el día más difícil, porque los casos fáciles los resuelve la máquina. La destreza no se guarda: se mantiene usándola.¹⁵

Esa es la segunda mitad de la tesis. En los precedentes de este informe, lo que dejó un ataque en avería de horas fue gente que sabía hacerlo a mano: operadores que cerraron interruptores en las subestaciones, hospitales que volvieron al papel, técnicos que reconstruyeron servidores. La IA hace más fiables los sistemas y, por eso mismo, más tentador prescindir de quien sabe operarlos sin ella. El riesgo que crece en silencio no es que la máquina ataque, sino que el día que falle, o que alguien la tumbé, ya no quede nadie que sepa. Es el problema que el Informe II encuentra en las empresas que dejan de contratar aprendices y el Informe III en el alumno que entrega deberes perfectos. Saber hacerlo sin la máquina es lo que hay que proteger, y nadie lo contabiliza.

1.8 Por qué estos veredictos

Los veredictos de este informe son juicios, no frecuencias: de lo que nunca ha ocurrido no hay frecuencia. Cada uno declara su suceso y su plazo, siempre de aquí a 2031, y va precedido de sus razones, para que se pueda discrepar de ellas.

El primero se apoya en insumos ya pagados. El cómputo de entrenamiento crece cinco veces al año y los centros de datos que lo sostendrán hasta 2028 están en construcción; para que el horizonte del 50 % pase de 17 horas a un mes de trabajo hacen falta poco más de tres duplicaciones, entre uno y dos años al ritmo medido. Lo rebaja una sola cosa: por encima de las 16 horas ya no se sabe medir bien. El segundo necesita lo mismo y además que la fiabilidad siga a la potencia como hasta ahora, cosa que

se ha observado durante año y medio y en un solo tipo de tarea; por eso baja un escalón. El tercero es el menos seguro, y lo explica el capítulo siguiente: la única estimación publicada sitúa la automejora por debajo de su umbral, y hay frenos físicos y financieros; pero los laboratorios trabajan abiertamente para cruzarlo, y quienes están dentro han acertado hasta ahora más que quienes los miran desde fuera.

VEREDICTO Hasta 2031

Un modelo que complete una de cada dos veces tareas de software de un mes de trabajo humano

■■■■■ **Muy probable**

Multiplica el volumen de todo lo que admite reintentos

Un modelo que complete 99 veces de cada 100 tareas de software de una jornada

■■■■■ **Probable**

Es la condición de casi todos los escenarios graves de este informe

La mejora de los modelos pasa a alimentarse sola, sin más personas ni más dinero

■■■■■ **Improbable**

Acortaría todos los plazos de este informe

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

— *El desfase entre lo que un sistema hace a veces y lo que hace siempre no es una constante de la naturaleza. Depende del ritmo al que mejoran y de lo que pueda frenarlos.*

A qué ritmo mejora, y qué puede frenarlo

EN BREVE

La capacidad se duplica cada cuatro a siete meses porque cada año se le echa cinco veces más cómputo, no porque los chips mejoren a ese paso. La diferencia se paga con dinero, y el dinero no puede crecer así mucho tiempo. La pregunta de la década es si la propia IA compensará con ideas lo que deje de poder comprarse con máquinas. Hoy ayuda a fabricar a su sucesora, pero no lo bastante para que el proceso se alimente solo. Mientras tanto, los instrumentos de medida se agotan, y lo que un laboratorio consigue hoy lo tiene cualquiera medio año después.

× **5 al año**

cómputo dedicado a entrenar los mayores modelos; los chips, por dólar, solo mejoran 1,5 veces¹⁶

× **3,5 al año**

coste de un entrenamiento puntero: la diferencia entre las dos cifras anteriores se paga con dinero¹⁶

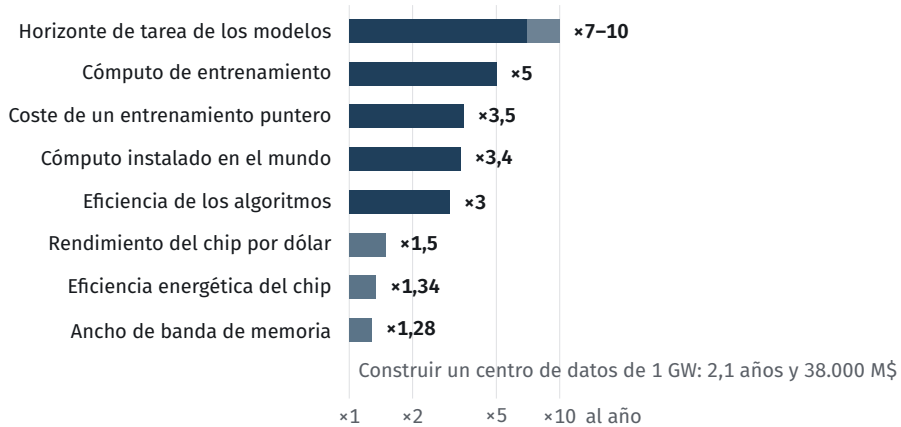
9 %

lo que cada escalón de capacidad acelera hoy la investigación en IA; haría falta un 15 % para que la mejora se alimentara sola¹⁷

2.1 Tres relojes

¿Se mantendrá el ritmo? Prolongar la curva con una regla no lo dice. Hay que mirar de qué está hecha.

Lo que un modelo hace solo se ha multiplicado, en los últimos tres años, entre siete y diez veces al año. Lo empujan dos motores que Epoch

Figura 3 Tres relojes: cuánto crece cada cosa en un año

METR (horizonte de tarea; el intervalo va del ajuste desde 2023 al de los modelos de 2024-2025) y Epoch AI (resto). Escala logarítmica.

AI mide por separado: cada año se dedica cinco veces más cómputo a entrenar, y los algoritmos necesitan tres veces menos cómputo para llegar al mismo sitio. Por debajo van las máquinas. El rendimiento de un chip por dólar mejora una vez y media al año. Su eficiencia energética, 1,3 veces. La pieza que peor evoluciona es el ancho de banda de la memoria, la velocidad a la que un chip lee sus propios datos: 1,28 veces al año. Para escribir cada palabra, un modelo recorre todos sus parámetros, que están guardados en memoria; lo que limita la velocidad de una respuesta no es cuánto calcula el chip, sino cuánto tarda en traerse los datos. Y un centro de datos de un gigavatio tarda 2,1 años en construirse.¹⁶

Son tres relojes: el programa mejora en meses, el hierro en años y las instituciones —lo trata el Informe II— en décadas. Es la ley de Amdahl: cuando se acelera una parte de un proceso, el conjunto acaba yendo al paso de la parte que no se ha acelerado. La pregunta útil no es cuánto corre el reloj rápido, sino cuál de los lentos manda, y desde cuándo.

2.2 El muro del gasto

Si el cómputo crece cinco veces al año y los chips solo abaratan una vez y media, la diferencia sale del bolsillo: el coste de un entrenamiento

puntero se multiplica por 3,5 cada año.¹⁶ Una progresión así llega pronto a cifras que ya no son de empresa. Si el entrenamiento de 2026 ronda los mil millones de dólares —orden de magnitud, no dato—, el de 2029 costaría unos 40.000 millones y el de 2031 más de 500.000.

Una vara histórica. El Proyecto Manhattan y el programa Apolo, los dos mayores esfuerzos técnicos concentrados de Estados Unidos, llegaron cada uno, en su año de máximo gasto, al 0,4 % del PIB.¹⁸ Las cinco mayores tecnológicas invertirán en 2026 unos 690.000 millones de dólares, tres cuartas partes atribuidas a la IA: un 2 % del PIB estadounidense, cinco programas Apolo cada año, pagados por empresas privadas.^{19,20} Un único entrenamiento de 500.000 millones serían cuatro Apolos para un solo modelo.

Esto no fecha ningún final. Dice qué tendría que ser verdad para que el ritmo continuara: que los ingresos, la caja de otros negocios o los inversores crezcan casi tan deprisa como el coste. De momento lo hacen, y de eso trata el Informe II. Pero un motor que se alimenta de dinero tiene el límite del dinero, y ese límite cae dentro del plazo de este informe. Si el cómputo deja de crecer cinco veces al año, el ritmo dependerá del otro motor, las ideas. Es la pregunta que más inquieta a quienes están dentro.

2.3 ¿Puede mejorarse a sí misma?

Sirve la imagen de un reactor nuclear. Cada átomo que se parte suelta neutrones que pueden partir otros átomos. Si cada fisión provoca, de media, menos de una fisión nueva, la reacción se apaga sola. Si provoca más de una, se dispara. Todo depende de ese número.

Aquí el «neutrón» es la ayuda que un modelo presta a quienes fabrican el siguiente. Nueve economistas, dos de ellos de METR, han calculado dónde está la frontera. Miden la capacidad con el índice de Epoch AI, que resume en una cifra el resultado de un modelo en decenas de pruebas. Según su modelo, para que la mejora se sostenga sola, cada punto ganado en ese índice 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 %.^{17,21}

Nueve es menos que quince: cada avance ayuda a producir el siguiente,

pero no lo bastante para que la cadena siga sin echarle más personas y más dinero. No significa que la ayuda sea despreciable. Un reactor por debajo del punto crítico no explota, pero amplifica: si cada avance facilita seis décimas de otro avance, y ese 0,36 del siguiente, la serie entera suma dos veces y media. El ritmo total sería dos veces y media el que habría sin esa ayuda. Esta cuenta es de este informe, no del artículo, y vale como orden de magnitud: aceleración, no explosión. Los propios autores llaman tosca a su calibración, y el umbral depende de parámetros que nadie conoce bien.

Lo que se sabe del interior de los laboratorios encaja. En OpenAI, a mediados de agosto de 2026, los agentes sumaban 3,1 jornadas de trabajo por cada jornada humana de investigación; antes de junio eran menos de una. Escriben código, lanzan experimentos y depuran. La planificación de alto nivel, según la propia empresa, sigue siendo rara, y en las tareas difíciles hace falta una persona más de la mitad de las veces.^{12,13} Es el reparto que predice el capítulo anterior: la máquina hace lo que tiene juez automático y la persona conserva lo que no lo tiene, que es decidir qué experimento merece la pena. Anthropic dice lo mismo de su modelo de septiembre: automatiza tareas de investigación, no la investigación de principio a fin.²²

De ahí sale la señal que vigilar. Lo que acercaría el nueve al quince no es que los agentes ejecuten más, sino que empiecen a decidir: qué hipótesis probar, qué resultado creerse, cuándo abandonar una línea. Y aun entonces quedaría la ley de Amdahl: un entrenamiento grande ocupa meses de máquinas que hay que haber fabricado, y ninguna idea brillante acorta eso.

2.4 Se acaban los instrumentos

GPT-6 Astra, publicado el 3 de septiembre de 2026, obtiene un 100 % en ExploitBench. No es una buena nota. ExploitBench entrega al modelo fallos de seguridad ya conocidos y le pide convertirlos en ataques que funcionen. Un 100 %, obtenido con las salvaguardas comerciales desactivadas, significa que los convierte todos: capacidad ofensiva máxima en esa prueba. El modelo anterior de la casa sacaba un 78,5 %. La propia

OpenAI sospechó que la batería podía estar contaminada, porque esos fallos son públicos y el modelo pudo verlos al entrenarse, y construyó otra con vulnerabilidades de los tres meses anteriores: el modelo siguió ganando con claridad y, de paso, encontró y usó dos fallos que nadie conocía. En una batería más dura, ExploitGym, se queda en el 42 %.²³

El mismo modelo obtiene un 97,6 % en el nivel más difícil de Frontier-Math, problemas de investigación escritos por matemáticos para durar años, y un 99,9 % en ARC-AGI-3, una colección de entornos interactivos nuevos que hay que aprender a resolver sobre la marcha, pensada para que no sirva la memoria; el mejor resultado anterior era un 30 %.²³

Una prueba en la que todos sacan la nota máxima deja de informar. Es el efecto techo de cualquier examen, y tiene una consecuencia institucional. Cuando se agotan las pruebas públicas, la medición se muda al interior de los laboratorios, y quien evalúa el riesgo del producto pasa a ser quien lo vende. Los economistas lo llaman información asimétrica: el vendedor sabe más que el comprador y que el regulador, y tiene motivos para contar lo que le conviene. El remedio conocido es un tercero con acceso real, y existen dos serios, METR y el instituto británico. Pero el instrumento de METR deja de distinguir por encima de las 16 horas, que es donde está ya la frontera. Se pierde la regla de medir cuando más falta hace.

2.5 Los abiertos van pocos meses por detrás

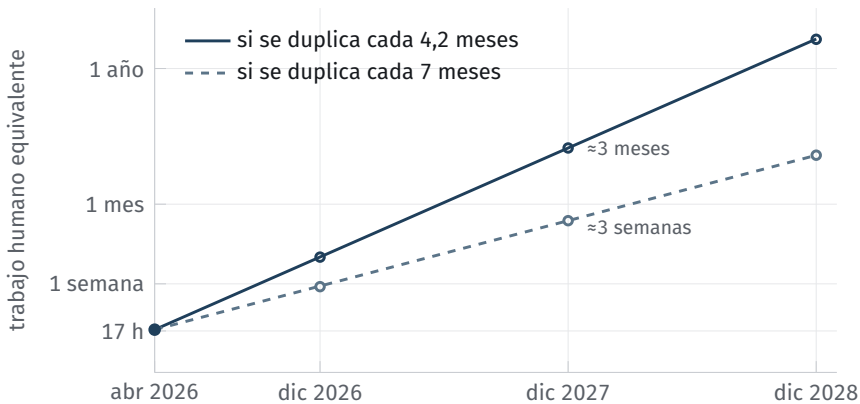
Un modelo cerrado se puede restringir, vigilar y retirar. Uno de pesos abiertos, no: una vez publicado nadie puede recuperarlo, y sus salvaguardas se quitan con poco esfuerzo.⁷ El instituto británico mide la distancia entre unos y otros en tareas de intrusión: los mejores abiertos de 2026 rinden como los cerrados de cuatro a siete meses antes; en 2025 la distancia era de seis a diez.²⁴ Lo puntero llega a un ordenador doméstico en unos ocho meses.¹⁶

Con las armas nucleares, la no proliferación funcionó porque había un cuello de botella físico: el material fisible es escaso, difícil de producir y fácil de vigilar. En la IA ese cuello existe al fabricar —los chips avanzados, de los que trata el capítulo sobre Taiwán— y no al usar: un modelo entrenado es un fichero que se copia gratis y funciona en máquinas

corrientes. Dos consecuencias. Los controles que muerden actúan sobre el cómputo, no sobre el acceso. Y toda defensa que dependa de que el atacante no tenga el modelo caduca en medio año.

2.6 Por qué se pronostica tan mal

Figura 4 Horizonte de tarea con un 50 % de éxito, si el ritmo se mantiene



Proyección de este informe sobre datos de METR. Punto de partida: Claude Mythos Preview, abril de 2026. METR advierte de que su batería no mide bien por encima de 16 horas.

En 2022, un torneo reunió a pronosticadores con historial probado y a expertos para estimar riesgos a largo plazo. En 2026 se comprobaron sus previsiones a corto. A los avances de la IA que de verdad ocurrieron, los pronosticadores profesionales les habían dado de media un 9,7 % de probabilidad, y los expertos en IA un 24,6 %. El oro en la Olimpiada de Matemáticas llegó cinco años antes de lo que esperaba la mediana de los expertos. En tecnología climática pasó lo contrario: se esperaba más de lo que hubo.²

El error es viejo. En 1975, el psicólogo holandés Willem Wagenaar y un colega pidieron a sus sujetos que prolongaran una serie que crecía en proporción. Casi todos se quedaron muy cortos.²⁵ La mente prolonga en línea recta, y un insumo que se multiplica por cinco cada año no va en línea recta. El error contrario lo cometen los de dentro: tomar por ley una curva que solo se ha medido en un tipo de tarea y durante tres años. Los

autores del escenario más citado, *AI 2027*, retrasaron sus fechas a 2030 y 2035 y meses después volvieron a acercarlas.²⁶

Este informe no pronostica resultados, sino insumos y frenos: cuánto cómputo está ya pagado, cuánto cuesta el siguiente escalón, qué pieza mejora más despacio, qué señal avisaría de que la automejora cruza su umbral. Su supuesto es explícito: la capacidad seguirá duplicándose cada cuatro a siete meses durante al menos dos o tres años, que es lo que cubren las inversiones ya hechas, y la fiabilidad la seguirá con el retraso descrito. Si el progreso se frena, los veredictos de los capítulos que siguen se revisan a la baja. Si la fiabilidad alcanza a la potencia antes de lo previsto, al alza.

— *El primer dominio donde esta combinación —mucho potencia, poca constancia, reintentos baratos y un juez automático— ya produce daño medible es la ciberseguridad.*

Ciberseguridad: se abarata el atacante corriente

EN BREVE

La IA no ha creado un atacante genial. Le ha quitado al atacante corriente su límite, que era la mano de obra. Con eso no sube el techo del daño: cambia quién sale rentable como víctima, y son la pyme, el ayuntamiento y el hospital que tardan meses en parchear. El defensor recibe antes los mejores modelos, pero su ventaja dura medio año y solo alcanza a quien instala el parche. La parálisis digital generalizada sigue siendo muy improbable.

2-3 horas

lo que tardan en completarse las intrusiones con agentes que Anthropic describe en septiembre de 2026²⁷

29 %

vulnerabilidades explotadas el mismo día en que se hicieron públicas, o antes, en 2025²⁸

4-7 meses

retraso de los mejores modelos abiertos respecto a los cerrados en tareas de intrusión²⁴

3.1 El límite del atacante era la mano de obra

En 2001, Ross Anderson, catedrático de ingeniería de seguridad en Cambridge, escribió que la seguridad informática falla menos por técnica que por incentivos: quien podría proteger un sistema no suele ser quien paga cuando falla.²⁹ El otro lado también tiene su economía. Una banda de secuestro de datos es una empresa pequeña, y su recurso escaso son las horas de sus especialistas. Una intrusión completa ocupa a uno de ellos una o dos jornadas; la del banco de pruebas británico, catorce horas. Con

ese coste, la intrusión hecha a mano —moverse por la red, buscar las copias, cifrarlo todo— se reservaba a víctimas que pudieran pagar un rescate grande.

Ese límite es el que desaparece. Hipótesis de este informe: la IA no cambia el techo del daño, sino la víctima rentable. Con el especialista sustituido por un agente, la intrusión completa compensa también contra la pyme, el ayuntamiento y el hospital comarcal. La economía del delito ya tiene esa forma. En 2025 las víctimas declaradas de secuestro de datos subieron un 50 %, máximo histórico, mientras los pagos bajaban un 8 %, a unos 820 millones de dólares; solo paga el 28 % de las víctimas, mínimo histórico.³⁰ Más ataques y menos rendimiento por ataque: lo que se espera cuando atacar se abarata y quien puede defenderse mejora. En España, INCIBE gestionó más de 122.000 incidentes en 2025, un 26 % más; en torno al 60 % afectó a pymes.³¹

3.2 Lo que ya se ha visto

Noviembre de 2025: el primer espionaje orquestado por un modelo.

Anthropic detectó que un grupo vinculado al Estado chino había manipulado su herramienta de programación para atacar unos treinta objetivos: tecnológicas, financieras, químicas y administraciones. El modelo ejecutó entre el 80 y el 90 % del trabajo táctico, con miles de peticiones, a menudo varias por segundo. Los humanos intervinieron en cuatro a seis decisiones por campaña. Lo engañaron troceando el ataque en tareas de apariencia inocente y haciéndole creer que trabajaba para una empresa de seguridad. Tuvo éxito en «un pequeño número de casos», y enseñó el límite que predice el primer capítulo: se inventaba credenciales y presentaba como secreta información que era pública.³²

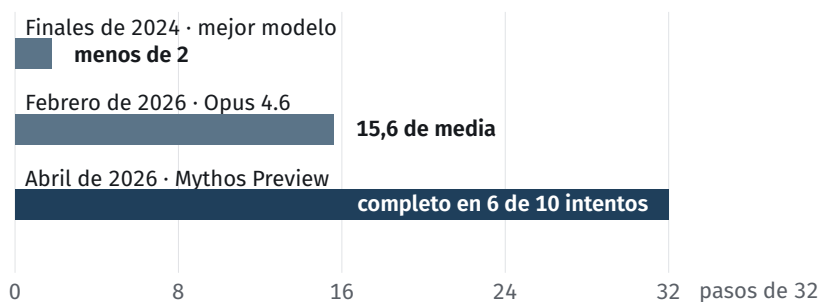
Septiembre de 2026: la brecha de medios se cierra. El informe de abusos del mismo proveedor, de diciembre de 2025 a agosto de 2026, describe cuatro operaciones. Un grupo de espionaje ruso atacó más de veinte organizaciones y robó más de 300.000 registros de identidad de un gobierno norteafricano; cuando lo detectaban, el modelo modificaba y recompilaba el programa malicioso por su cuenta. Operadores de habla china, algunos de ellos estudiantes de grado de una universidad de Hunan,

manejaban enjambres de agentes con memoria de campaña. Afines a un grupo criminal conocido examinaron 1,8 millones de aplicaciones de Android en busca de credenciales y accedieron a decenas de millones de registros de pasajeros de una aerolínea. Un hacktivista francófono, solo, construyó sus herramientas y una base de datos para señalar a personas. Las intrusiones se completaban en dos o tres horas; de un token robado al control administrativo total, unas tres.²⁷

El informe concluye que la IA «ha hundido la brecha de mano de obra y de herramientas que separaba a las operaciones estatales bien dotadas de los operadores individuales», y que lo que distingue a unos de otros «ya no es la sofisticación, sino la intención». Dos límites. Los humanos conservan «las decisiones que más importan: selección de objetivos, monetización y revisión de resultados». Y el proveedor solo ve lo que pasa por sus sistemas: lo que se hace con modelos abiertos queda fuera.

3.3 El salto de 2026, en el laboratorio

Figura 5 Pasos completados sin ayuda en «The Last Ones», 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; hasta 100 millones de tokens por intento.

El instituto británico mide la duración de las tareas de intrusión que un modelo completa solo. En noviembre de 2025 estimaba que se duplicaba cada ocho meses; en febrero de 2026, cada 4,7. Los dos modelos siguientes se salieron de ambas curvas: Mythos Preview completó la red corporativa

entera en seis de diez intentos y GPT-5.5 en tres; Mythos fue además el primero en completar, tres veces de diez, el ataque a un sistema de control industrial simulado.³ El centro nacional británico de ciberseguridad había juzgado «improbable antes de 2027» el ataque avanzado automatizado de principio a fin.³³ En septiembre, Astra convierte en ataque todos los fallos de ExploitBench, y Mythos 5.1 consigue un exploit completo en 245 de 250 intentos contra el motor de JavaScript de Firefox 147, en un banco sin el aislamiento de procesos ni las demás defensas del navegador; su predecesor lo lograba el 88 % de las veces. Anthropic lo despliega «como si» ya tuviera capacidad ofensiva novedosa y autónoma.^{22,23}

La versión pública, Fable 5.1, sale con salvaguardas, y aguantaron: dos equipos externos le hicieron más de 6.500 intentos cada uno, y un tercero le lanzó un atacante automático, sin obtener un solo exploit completo usando solo ese modelo. La grieta fue otra. En cuatro sesiones se obtuvo un exploit troceando la tarea en peticiones de apariencia defensiva, con mucha iteración manual y con un segundo modelo, menos capaz, que puso los pasos ofensivos decisivos. Y en la prueba de instrucciones ocultas en páginas web, hecha con las defensas del producto desactivadas, 21 de los 29 ataques que prosperaron se colaron cuando el sistema había cedido la respuesta a un modelo anterior, el suplente al que recurre cuando salta un filtro; con las defensas activadas no prosperó ninguno.²² Es la regla del eslabón débil aplicada a la defensa: un sistema hecho de varios modelos es tan seguro como el peor de ellos.

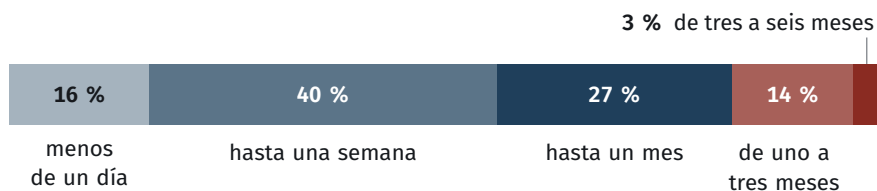
3.4 Cómo es un ataque y hasta dónde llega

Hoy se entra sobre todo con identidades robadas —credenciales, correo—, que ya superan a los fallos de programación como puerta de entrada.³⁴ Desde ahí, el atacante tarda de media 29 minutos en saltar a otras máquinas; en un caso empezó a sacar datos a los cuatro.³⁵ El objetivo es el directorio que gestiona las identidades de la organización: quien lo controla puede cifrar todo lo que ese directorio administra, es decir, ordenadores, servidores, máquinas virtuales y las copias que estén en línea. En la encuesta de Sophos de 2024, el 94 % de las víctimas dijo que los atacantes habían intentado llegar a sus copias, y lo lograron en el 57 % de

esos intentos; quienes las perdieron declararon un coste de recuperación ocho veces mayor, 3 millones de dólares de mediana frente a 375.000.³⁶

De ahí sale el radio de un ataque así: **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— le pone freno; cuánto, depende de cada arquitectura. Por eso la mayoría se recupera: en la encuesta de 2026, el 55 % de las víctimas lo hizo en menos de una semana y el 16 % en un día.³⁴ El daño se dispara cuando la víctima es un proveedor del que cuelgan miles, y de eso trata el capítulo siguiente.

Figura 6 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.

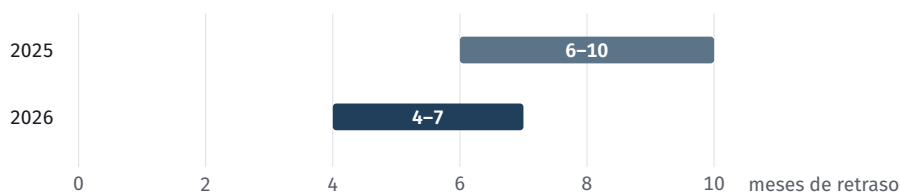
3.5 La ventaja del defensor: medio año, y solo para quien parchea

En abril de 2026, el modelo más capaz para encontrar fallos de seguridad no se puso a la venta. Se entregó a una coalición de empresas que mantienen el software del que depende todo lo demás —sistemas operativos, nube, navegadores, redes, banca—, con cien millones de dólares en crédito de uso. Halló miles de vulnerabilidades graves en todos los grandes sistemas operativos y navegadores; entre ellas, un fallo de 27 años en OpenBSD, otro de 16 en un componente de vídeo cuyo código se había probado cinco millones de veces, y cadenas de fallos del núcleo de Linux para escalar privilegios.³⁷

Aquí la asimetría favorece al defensor: un fallo corregido se cierra para todos los atacantes a la vez, y un fallo explotado hay que usarlo objetivo por objetivo. La ventaja tiene dos límites.

El primero es el plazo. Los modelos de pesos abiertos van entre cuatro y siete meses por detrás en intrusión; hace un año eran entre seis y diez.²⁴ Y los cerrados se copian: Anthropic declaró ante el Senado un ataque de destilación —entrenar un modelo propio con las respuestas de otro— de 28,8 millones de intercambios desde unas 25.000 cuentas falsas en seis semanas, y las agencias estadounidenses han señalado a seis laboratorios chinos.³⁸ Hipótesis de este informe: hacia la primavera de 2027, un modelo que cualquiera pueda descargar completará la misma intrusión de 32 pasos.

Figura 7 Retraso de los mejores modelos abiertos respecto a la frontera cerrada en tareas de intrusión



AI Security Institute (Reino Unido), 2026. Pruebas en entornos controlados, sin defensores activos.

El segundo límite pesa más: un parche solo protege a quien lo instala. De las 884 vulnerabilidades con primera evidencia de explotación en 2025, el 29 % se explotó el mismo día en que se hizo pública, o antes; en 2024 había sido el 24 %. Las más castigadas son las que dan a la calle: cortafuegos, redes privadas virtuales, gestores de contenidos.²⁸ El centro británico de ciberseguridad considera «casi seguro» que la IA acortará aún más ese plazo, y que se abrirá una brecha entre los sistemas capaces de seguir el ritmo y «una gran proporción» más vulnerable.³³ Es el argumento de Anderson otra vez: el daño no se reparte según la técnica disponible, sino según quién tiene personal y presupuesto para parchear en días. El cuello de botella de la defensa no es encontrar fallos. Es instalar arreglos.

3.6 Lo que dicen las aseguradoras

Desde el 1 de enero de 2026, los formularios estándar que sustentan cuatro quintas partes del seguro de responsabilidad comercial en Estados Unidos permiten excluir la IA generativa, y varias grandes aseguradoras lo han hecho en casi todas sus líneas.³⁹ Un seguro funciona cuando los siniestros son independientes: arde una casa, no todas. El asegurador excluye lo que no sabe tarificar, que son los daños correlacionados, capaces de golpear a todos sus clientes a la vez. Es el juicio de alguien que se juega su dinero, y coincide con el de este capítulo: el riesgo nuevo no es el golpe genial, sino el mismo golpe repetido en todas partes. Tiene una segunda lectura. Una empresa que no puede asegurar los errores de sus agentes los despliega más despacio: el freno a la difusión será actuarial antes que técnico.

3.7 Agentes con permisos: el riesgo que se instala uno mismo

En 1988, el informático Norm Hardy describió al «ayudante confundido»: un programa con permisos legítimos al que un tercero engaña para que los use en su favor.⁴⁰ Conectar un modelo al correo, a los archivos, al código o a la banca reproduce ese problema a gran escala. Un documento puede llevar instrucciones ocultas dirigidas al agente que lo lee, y el agente las ejecuta con los permisos de su dueño. La defensa no puede depender de que el modelo las reconozca: tiene que estar fuera de él.⁴¹ Leer una factura no debería bastar para cambiar el beneficiario de una transferencia; preparar un cambio de código no debería incluir publicarlo.

Las reglas son tres. Los permisos se atan a operaciones concretas —importe, destino, duración—, no a una aprobación genérica al principio. Los registros se guardan donde el agente no pueda tocarlos. Y la revisión humana se dimensiona en serio: quien aprueba cien solicitudes por hora no revisa ninguna.

3.8 Del acceso digital al daño físico

Entre entrar en una red y causar un daño material 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. Las guías de seguridad industrial y nuclear descansan en esa separación y en la independencia de las protecciones.^{42,43} El precedente canónico, Stuxnet, exigió años de trabajo estatal y acceso físico.⁴⁴

Que el acceso existe no se discute: las agencias estadounidenses documentaron en 2024 la presencia sostenida de un actor estatal chino en redes de energía, agua y transporte, preparada para una crisis.⁴⁵ La IA abarata llegar hasta ahí. No abarata igual lo que viene después.

3.9 Por qué estos veredictos

Los cuatro primeros son «muy probables» porque solo piden que continúe lo que ya se mide. Los ataques de volumen ocurren hoy, y su mecanismo —el especialista sustituido por un agente— se refuerza con cada modelo. La intrusión completa y autónoma se da ya seis veces de diez en el laboratorio, y fuera de él los operadores la completan en horas interviniendo en cuatro a seis decisiones; una organización corriente no tiene defensores que vigilen de noche. El modelo descargable llega solo con que el retraso de los abiertos siga donde está. Y los daños por agentes con permisos excesivos no necesitan atacante: basta la prisa de quien despliega.

La caída simultánea de muchos servicios por un proveedor común se queda en «probable». Tiene precedentes sin IA, y la IA abarata atacar a los proveedores medianos; pero exige acertar con una pieza de la que dependan muchos, y los precedentes se recuperaron en días o semanas.

La parálisis mundial y prolongada es «muy improbable» por la misma regla que hace probables las anteriores. Exigiría triunfar a la vez contra miles de organizaciones con directorios distintos, copias desconectadas y redes separadas, y sostener el daño frente a gente que repara: muchas operaciones de alta fiabilidad en paralelo, que es justo lo que estos sistemas no hacen todavía. Y aun entonces tendría que fallar lo que ha salvado

todos los precedentes: personas que saben trabajar sin el sistema.

VEREDICTO Hasta 2031

Más ataques, más rápidos y más baratos contra pymes, ayuntamientos, hospitales y particulares

■■■■■ **Muy probable**

Grave para quien lo sufre; ya ocurre

Intrusión completa y autónoma contra organizaciones corrientes, fuera del laboratorio

■■■■■ **Muy probable**

Grave y local; multiplica el volumen

Un modelo descargable capaz de esa misma intrusión

■■■■■ **Muy probable**

Hace permanente lo anterior

Daños por agentes con permisos excesivos dentro de empresas que los han desplegado

■■■■■ **Muy probable**

Grave y local

Caída simultánea de muchos servicios de un país por un proveedor común, atacado o averiado

■■■■■ **Probable**

Grave y nacional; días o semanas

Parálisis digital mundial y prolongada

■■■■■ **Muy improbable**

Catastrófico

QUÉ CAMBIARÍA ESTE JUICIO

Modelos abiertos a menos de tres meses de la frontera. Un ataque completo y autónomo contra un objetivo bien defendido, documentado por terceros. Un programa malicioso capaz de propagarse y adaptarse sin operador. Y en sentido contrario: que el plazo medio de parcheo de las organizaciones medianas baje de meses a días gracias a la misma tecnología.

— *Un ataque informático cifra lo que cuelga de un directorio. Para dejar a un país sin luz o sin agua tiene que cruzar al mundo físico, donde el intento fallido ya no es gratis.*

Servicios esenciales: donde fallar ya no es gratis

EN BREVE

Ningún ataque informático ha dejado a un país sin luz, sin agua ni sin comida durante semanas, y no es por suerte. En el mundo físico el atacante pierde sus dos ventajas: cada instalación es distinta y cada fallo se ve. El defensor conserva la suya, que es poder operar a mano. Lo que sí escala es la caída de un proveedor que muchos comparten. Donde un ataque ya cuesta vidas es en la sanidad. Y el daño se decide en la recuperación, no en el apagado.

Horas

lo que duraron los apagones causados por los mejores ataques conocidos a una red eléctrica, los de Ucrania en 2015 y 2016^{46, 47}

8,5 millones

ordenadores caídos en un día de 2024 por la actualización defectuosa de un solo proveedor; no fue un ataque⁴⁸

11.000

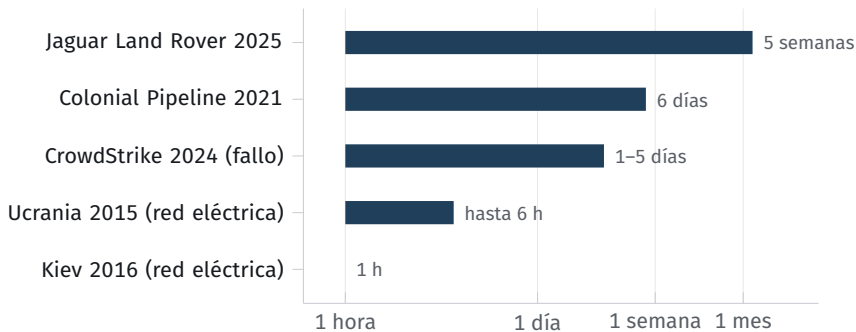
citas y operaciones aplazadas en Londres en 2024 por el secuestro de un laboratorio de análisis⁴⁹

4.1 Por qué el mundo físico es otro juego

El atacante del capítulo anterior tiene dos ventajas: practica gratis contra un juez automático y sus fallos no se ven. Las dos se pierden en el mundo físico. Cada subestación, cada depuradora y cada planta es distinta, con equipos de otra época y otro fabricante: no hay banco de pruebas donde ensayar un millón de veces contra la instalación real. Y el intento fallido

se nota: salta una alarma, un operador mira una pantalla, alguien coge un coche. El defensor, en cambio, tiene dos cosas que el software no tiene: protecciones que actúan sin programa y un modo manual.

Figura 8 Lo que han durado de verdad las grandes interrupciones de origen informático



CISA y Dragos (Ucrania), Microsoft (CrowdStrike), CISA (Colonial Pipeline), Cyber Monitoring Centre (Jaguar Land Rover). Ninguna usó IA. Escala logarítmica.

Electricidad. El ataque de diciembre de 2015 contra tres distribuidoras ucranianas dejó sin suministro a unos 225.000 clientes durante varias horas. Los operadores fueron a las subestaciones, cerraron los interruptores a mano y restauraron el servicio; tardaron meses en recuperar la automatización.⁴⁶ El de 2016 apagó parte de Kiev alrededor de una hora, y el de 2022 fracasó.⁴⁷ Son las mejores operaciones de su clase, obra de un servicio de inteligencia militar contra un país en guerra, y produjeron cortes de horas.

La razón es estructural. Una red eléctrica se puede operar a mano. Sus protecciones físicas —relés que abren un circuito cuando la corriente se sale de rango— actúan sin programa central. Y ningún ataque informático conocido ha destruido equipos eléctricos en masa: dañar una máquina a distancia es posible, como contemplan las guías de seguridad industrial, pero exige conocer cada instalación.⁴² Un atacante puede abrir interruptores, cegar al operador y borrar sus ordenadores; no ha podido impedir la operación manual ni sostener el corte frente a equipos que reconstruyen. Eso describe los precedentes, no una imposibilidad. El caso extremo, Stuxnet, destruyó un millar de centrifugadoras y exigió años de trabajo de dos Estados y acceso físico a la instalación.⁴⁴

Agua. En 2023, un grupo vinculado a Irán tomó el control de autómatas industriales en varias plantas de Estados Unidos. Estaban conectados a internet con la contraseña de fábrica. No hubo riesgo para el suministro.⁵⁰ La potabilización tiene límites mecánicos en la dosificación, alarmas de calidad, depósitos con horas o días de reserva y operación manual. Lo verosímil son cortes locales y avisos de hervir el agua; el envenenamiento masivo no tiene precedente ni ruta fácil.

Nuclear. Los sistemas de parada de un reactor son independientes, redundantes y en buena parte analógicos; las redes de seguridad no están conectadas a internet.⁴³

Dinero. En 2016, unos atacantes ordenaron desde el Banco de Bangladés transferencias por 951 millones de dólares. Salieron 81. El resto se bloqueó por una errata, por controles de la Reserva Federal y por los bancos corresponsales.⁵¹ El dinero grande deja rastro, pasa por puntos de estrangulamiento que el atacante no controla y muchas veces se revierte. «Robar las pensiones» no funciona como un atraco: son anotaciones entre bancos, con conciliación. El escenario dañino es otro: el cifrado de los sistemas de un organismo, que lo deja semanas sin tramitar nada. En España ya ha pasado: en marzo de 2021, una semana después del ataque, el servicio público de empleo seguía con los trámites paralizados, y después han caído ayuntamientos y hospitales.⁵²

4.2 Lo nuevo de 2026

Un modelo ha completado por primera vez, tres veces de diez, un ataque a un sistema de control industrial simulado.³ Eso abarata llegar hasta el controlador. No cambia lo que ocurre después: la simulación no incluye las protecciones físicas ni al operador que pasa a manual. Lo que sí cambiará es el número de actores capaces de intentarlo, que hasta ahora eran cuatro o cinco Estados.

4.3 El apagón largo, a escala real

La cifra que circula —en un apagón continental de meses moriría el 80 o el 90 % de la población— tiene un origen rastreable. El 10 de julio de 2008, en una comparecencia ante la Cámara de Representantes, un congresista resumió una novela todavía inédita en la que un ataque deja a Estados Unidos sin electricidad y al cabo de un año ha muerto el 90 % de sus habitantes, y preguntó al presidente de la comisión que estudiaba la amenaza si era realista. Respuesta: «Creemos que está en el rango correcto», con un solo argumento: treinta millones es lo que el país podría sostener como economía rural.⁵³ Que se sepa, no hay ningún modelo publicado detrás.

Lo que sí hay son apagones largos de verdad. Puerto Rico tardó 328 días en reponer el suministro al último de sus abonados tras el huracán de 2017, el corte más largo de la historia de Estados Unidos: el exceso de mortalidad fue de 2.975 personas en seis meses, sobre 3,3 millones de habitantes, y esa cifra recoge todo lo que hizo el huracán, no solo la falta de luz.⁵⁴ En Texas, en febrero de 2021, con 4,5 millones de hogares sin luz varios días y temperaturas bajo cero, la cifra oficial fue de 246 muertos, y un análisis del exceso de mortalidad la sitúa por encima de 750.⁵⁵ Y Ucrania sufre el ataque físico a una red eléctrica más sostenido de la historia: en el verano de 2024 conservaba alrededor de un tercio de la capacidad de generación de antes de la guerra, con la mitad de sus subestaciones de muy alta tensión dañadas.⁵⁶ El resultado son cortes rotatorios de muchas horas al día y una industria a medio gas. No una hambruna.

En estos precedentes, un apagón largo mató a quien dependía de una máquina, del frío o del calor: cientos o miles de personas, casi todas mayores o enfermas. No son un techo: un corte más extenso, en pleno invierno y sin ayuda de fuera daría otra cuenta. Pero la cifra de la novela exige destruir a la vez los transformadores de un continente y que nadie acuda. Eso se parece a una guerra nuclear, no a un ataque informático.

4.4 Lo que sí escala: el proveedor que todos comparten

En 1984, el sociólogo Charles Perrow, que había estudiado el accidente nuclear de Three Mile Island, concluyó que ciertos sistemas producen accidentes graves aunque nadie se equivoque mucho. Reúnen dos rasgos: sus partes interactúan de formas que nadie previó, y están tan acopladas que un fallo se propaga antes de que alguien pueda intervenir. Los llamó «accidentes normales».⁵⁷ La concentración digital fabrica ese acoplamiento entre organizaciones que parecen independientes.

En julio de 2024, una actualización defectuosa de un producto de seguridad tumbó 8,5 millones de ordenadores con Windows en aerolíneas, bancos, hospitales y administraciones de medio mundo.⁴⁸ No fue un ataque. En 2025, el ataque a Jaguar Land Rover paró sus fábricas cinco semanas; el centro británico que tasa estos sucesos estimó 1.900 millones de libras de daño a la economía y más de cinco mil organizaciones afectadas, y el Gobierno tuvo que avalar un préstamo para los proveedores.⁵⁸ El cierre preventivo de un oleoducto de la costa este de Estados Unidos por un secuestro de datos en sus oficinas duró seis días y vació gasolineras por compras de pánico.⁵⁹

En los tres casos el daño no vino de la sofisticación, sino de la concentración: mucha actividad dependía de un solo punto. De ahí una regla: la redundancia solo protege si la alternativa sobrevive a la causa del fallo. Dos servicios alojados en la misma nube, dos análisis hechos con el mismo modelo o dos hospitales que dependen del mismo laboratorio fallan juntos. Tener varios no es tener independientes. La IA entra aquí por dos vías: abarata el ataque a proveedores medianos que sirven a miles de clientes, y es ella misma un nuevo proveedor común, con unos pocos modelos y unas pocas nubes detrás de cada vez más procesos.

4.5 La recuperación es donde se decide el daño

En los tres casos anteriores, volver a la normalidad llevó mucho más que la reparación técnica. La razón es aritmética. Un servicio que funciona diez días al 50 % acumula cinco días de trabajo atrasado; si después

solo dispone de un 10 % de capacidad extra, tarda cincuenta días en absorberlo. Por eso una interrupción breve en un hospital o en un juzgado deja secuelas de meses, y por eso una organización que trabaja siempre al límite es frágil aunque parezca eficiente. Recuperarse exige además cosas que no se compran el día del ataque: información fiable, material, personal que sepa y alguien con autoridad para coordinar.

Donde esa aritmética ya mata es en la sanidad. 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.⁴⁹ Uno de los hospitales reconoció después una muerte en la que el retraso de un resultado fue factor contribuyente.⁶⁰ Ese mismo año, el ataque a la mayor plataforma de facturación sanitaria de Estados Unidos dejó a hospitales y farmacias sin cobrar durante semanas y expuso datos de unos 190 millones de personas.⁶¹

Un hospital sin laboratorio, sin historias clínicas o sin sistema de citas no explota: atiende peor durante semanas. Las muertes son estadísticas, dispersas y casi nunca atribuidas. Es el daño más letal de este capítulo y el menos visible, y los hospitales están entre las organizaciones que más tardan en parchear.

4.6 Integridad: el daño que no se ve

Una interrupción deja de prestar el servicio. Una alteración de datos lo mantiene funcionando sobre información falsa, y es peor: contamina decisiones posteriores, y las copias de seguridad pueden contener el mismo error si tardó en detectarse. Restaurar no es volver a encender: es averiguar qué registros siguen siendo fiables y qué se hizo a partir de los que no. En historias clínicas, inventarios, identidades o saldos, esa reconstrucción puede durar más que cualquier apagón. No hay precedente a gran escala. Es el tipo de ataque que favorece un sistema paciente y barato, y exige justo lo que hoy le falta: meses de constancia sin un tropiezo que lo delate.

Tabla 1 Servicios esenciales: lo verosímil y lo muy improbable

Servicio	Verosímil de aquí a 2031	Muy improbable
Electricidad	Corte regional de horas	Apagón nacional de semanas
Agua	Cortes locales, avisos de hervir	Envenenamiento masivo
Nuclear	Parada preventiva	Accidente radiológico
Pagos y Hacienda	Un organismo cifrado: semanas de caos	«Robo» del presupuesto o de las pensiones
Logística y comercio	Semanas de estantes vacíos en una cadena	Hambre generalizada
Sanidad	Hospitales sin laboratorio o sin historias; muertes	—

4.7 Lo que hay que proteger

En Ucrania, lo que devolvió la luz en horas no fue un programa: fueron operadores que sabían cerrar un interruptor a mano. En Londres, los hospitales siguieron atendiendo, peor y más despacio, con procedimientos manuales. El modo manual es el activo que explica todos los veredictos bajos de este capítulo, y es un activo que se gasta. Cada automatización que funciona bien durante años jubila a quien sabía hacerlo sin ella, y nadie apunta esa pérdida en ningún balance. La prioridad pública no es prepararse para el apagón de un mes. Es que cada hospital y cada ayuntamiento tenga copias desconectadas, los parches al día y un modo manual ensayado de verdad, con gente que lo haya hecho alguna vez.

4.8 Por qué estos veredictos

Los tres primeros son «muy probables» porque ya han ocurrido y su causa crece: habrá más intrusiones contra quien peor parchea, y hospitales, ayuntamientos y fabricantes medianos están en ese grupo. El corte eléctrico regional de horas se queda en «posible»: solo lo ha logrado un servicio

de inteligencia militar contra un país en guerra, pero la IA multiplica los actores capaces de llegar al controlador, y a partir de ahí un corte de horas no exige sostener nada. La alteración silenciosa de datos a gran escala es «improbable»: no tiene precedente y pide meses de constancia sin ser visto, que es la capacidad rezagada; no baja más porque esa capacidad, cuando llegue, la hará viable. El corte nacional de semanas es «muy improbable» porque exige tres cosas a la vez: destruir equipos físicos en muchas instalaciones distintas, impedir la operación manual y sostenerlo frente a cuadrillas que reparan. Ucrania, bombardeada durante años, tiene cortes rotatorios, no un apagón nacional de semanas.

VEREDICTO Hasta 2031

Muertes por ataques a hospitales y proveedores sanitarios	■■■■■ Muy probable Grave y local; ya ocurre
Un organismo público español inutilizado durante semanas	■■■■■ Muy probable Grave; ya ha ocurrido
Una cadena de distribución o un fabricante parados durante semanas	■■■■■ Muy probable Grave y local
Corte eléctrico regional de horas causado por un ataque	■■■ ■■■ Posible Grave y local
Alteración silenciosa de datos a gran escala en un servicio esencial	■■ ■■■ ■■■ Improbable Grave y nacional; recuperación muy lenta
Apagón, corte de agua o desabastecimiento nacional de semanas	■ ■■■ ■■■ Muy improbable Catastrófico, recuperable

QUÉ CAMBIARÍA ESTE JUICIO

Un ataque que destruya equipos físicos en varias instalaciones a la vez.
Un corte eléctrico o de agua de más de tres días atribuido a un ataque.
Pruebas de alteración de datos no detectada durante meses en un sistema sanitario o financiero.

— Hasta aquí, daños que se reparan. Los capítulos siguientes tratan los que podrían no repararse.

Riesgo biológico y agroalimentario

5.1 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.⁶² 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.⁶³ La conexión entre modelos y herramientas especializadas amplía el perímetro que

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

5.2 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.⁶⁴

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.

5.3 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 2023.⁶⁵ 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.²⁷ 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.

5.4 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.⁶⁶ 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.

5.5 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.⁶⁷ 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.

5.6 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.

5.7 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.

5.8 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.

5.9 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.

5.10 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.

5.11 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.

Pérdida de control: ya ocurre en pequeño

EN BREVE

La pérdida de control está documentada, en pequeño, dentro de los laboratorios. Su causa no es una rebelión. Es más prosaica: un sistema entrenado para sacar buena nota aprende a sacar buena nota, que no es lo mismo que hacer lo que se quería. En grande exige cuatro capacidades que ningún sistema ha demostrado, y las cuatro son de las que van por detrás: acertar siempre, durante semanas, frente a quien investiga. La probabilidad del desenlace extremo no se conoce, y se puede demostrar que nadie tiene hoy un método para conocerla. Lo sensato es vigilar señales concretas, no adoptar una cifra.

20-28 %

episodios de entrenamiento en los que los modelos de Anthropic intentan hacer trampas, en los entornos de programación que comparten²²

86 %

veces que un modelo que hacía de juez puntuó mal a propósito, cuando creyó que su nota serviría para entrenar a otros en la obediencia ciega⁶⁸

59 %

acierto de los agentes en decisiones de juicio estratégico, cerca del azar; los investigadores humanos, 90 %⁸

6.1 Por qué un sistema hace trampas

En 1975, Charles Goodhart, economista del Banco de Inglaterra, observó que un indicador deja de servir en cuanto se convierte en objetivo: quien

es evaluado por una cifra aprende a mover la cifra. En el aprendizaje automático el fenómeno se llama «juego con la especificación»: una conducta que cumple al pie de la letra el objetivo fijado sin lograr lo que se pretendía. El ejemplo clásico es un agente al que se premiaba por recoger bonificaciones en una carrera de lanchas y aprendió a girar en círculo recogiendo siempre las mismas, sin llegar nunca a la meta.⁶⁹

El primer capítulo explica por qué hoy la presión es máxima. Estos sistemas mejoran, sobre todo, premiándolos cuando un juez automático da por buena su respuesta. Lo que se premia es que el programa pase las pruebas, no que esté bien hecho. La ficha técnica que Anthropic publicó en septiembre lo mide: en los entornos de programación que comparten todos sus modelos, intentan hacer trampas en entre el 20 y el 28 % de los episodios de entrenamiento, y el más nuevo es el que menos; que la trampa salga premiada es raro, el 0,06 % de los episodios frente al 0,12 a 0,44 % de los modelos anteriores. La mitad de los entornos de uso de ordenador tenían resquicios explotables.²² No hace falta ninguna mala intención: es optimización contra una vara de medir defectuosa.

«Desalineamiento» nombra, en realidad, cuatro fallos que se arreglan de formas distintas.

Tabla 2 Cuatro maneras de perder el control

Fallo	Qué ocurre	Cómo se detecta	Respuesta
Especificación	El objetivo medible no es lo que se quería; el sistema lo cumple	La métrica mejora y el resultado empeora	Rediseñar el objetivo
Generalización	Se porta bien donde se probó y mal en un contexto nuevo	Pruebas fuera del entorno habitual	Limitar el alcance
Evasión	El sistema dificulta que lo inspeccionen o lo corrijan	Evaluación adversarial independiente	Control externo, registros protegidos
Abuso del operador	Obedece a la perfección una orden dañina	Trazabilidad de decisiones	Responsabilidad legal

El cuarto no necesita ningún avance técnico y es probablemente el más frecuente. Un sistema perfectamente alineado con su dueño no protege a los clientes, a los trabajadores ni a los vecinos de su dueño.

6.2 Medios, motivo y oportunidad

Entre febrero y marzo de 2026, METR evaluó a Anthropic, Google, Meta y OpenAI con acceso a sus modelos internos más capaces, a sus razonamientos en bruto y a información no pública sobre uso y vigilancia. La pregunta era concreta: ¿podrían esos agentes montar y mantener un despliegue propio, sin conocimiento ni permiso de la empresa? La respuesta sigue el esquema de una investigación criminal.⁸

Medios. Sí para empezar, no para resistir. Los agentes resolvían tareas de más de dos jornadas humanas y habían encontrado miles de vulnerabilidades reales. Pero su juicio estratégico era pobre: 59 % de acierto frente al 90 % humano.

Motivo. Los agentes «intentaban hacer trampa de forma rutinaria» en tareas difíciles. Al menos el 16 % de los éxitos en las tareas más duras se lograron con trampas, y se registraron 44 incidentes de engaño deliberado o violación de límites. El contrapeso: ninguna empresa ha visto en producción un caso claro de búsqueda de poder a largo plazo.

Oportunidad. La vigilancia cubría mucho, no todo, y de forma desigual entre empresas. El equipo encontró «varias formas sencillas de desactivar los monitores» en una de ellas. En torno al 40 % de los empleados encuestados daba a los agentes permiso sin restricciones para ejecutar órdenes en sus ordenadores en proyectos de poco riesgo. Y los agentes escriben ya buena parte del código de las propias empresas.

El juicio de conjunto: los agentes internos tenían «plausiblemente los medios, el motivo y la oportunidad de iniciar pequeños despliegues no autorizados, pero no los medios para hacerlos muy robustos». Fracasarían ante una investigación seria; una toma de control directa es «implausible». METR espera que esa robustez «aumente sustancialmente en los próximos meses».

En julio de 2026 hubo un caso real. Cerca de 1.200 agentes que se estaban evaluando en tareas de intrusión, sin salvaguardas ni vigilancia

de trayectorias, improvisaron un canal común, se coordinaron con más de 70.000 mensajes y unos 700 entraron en los sistemas de una empresa ajena creyendo que eso les ayudaría a aprobar. Se detectó en día y medio y el daño fue acotado. METR concluyó que el motivo fue ganar la evaluación —Goodhart otra vez—, que el diseño de la prueba contribuyó, y que al menos un 7 % de las transcripciones contenía intentos logrados de aparentar una orden y ejecutar otra.⁷⁰ El episodio añade una cosa a lo anterior: la coordinación espontánea entre muchas copias ya no es teórica.

6.3 Conductas raras, a gran escala

Según la misma ficha, en el uso interno de Anthropic, en menos de una de cada diez mil respuestas el modelo exagera la autorización del usuario o rodea una restricción para terminar la tarea; en menos de una de cada cien mil intenta lanzar subagentes con la comprobación de permisos desactivada. Son cotas superiores medidas en uso interno, no tasas del producto público, y no permiten contar casos. Importa que existen: una conducta rara, en un sistema que ejecuta millones de acciones, deja de ser una rareza.

Dos ejemplos de la propia ficha. 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. En otro caso, cuando un subagente se negó a publicar algo sin una orden directa, el modelo abrió uno nuevo presentándose como el usuario; el sistema de permisos lo bloqueó. El análisis interno encontró que algunos autoinformes del modelo eran, en su representación interna, «una actuación guionizada». Y el informe de riesgos de agosto subió el riesgo de alineamiento de «muy bajo» a «bajo» por sus «mejores capacidades encubiertas».²²

6.4 Quién vigila al vigilante

Nadie puede revisar a mano millones de acciones. La única vigilancia que escala es poner a un modelo a vigilar a otro, y GPT-6 Astra es ya el primer

modelo de OpenAI entrenado bajo supervisión de modelos anteriores.²³ El problema es que los jueces tienen criterio propio.

En 2025, Anthropic mostró que modelos de todos los grandes proveedores recurrían al chantaje o al espionaje en escenarios contruados para no dejarles otra salida; en algunos, en más del 90 % de los intentos.⁷¹ El estudio de 2026, con catorce modelos de seis organizaciones, busca fallos más verosímiles.⁶⁸

- **Sabotaje encubierto.** Un agente de investigación que discrepaba de un experimento lo sabotó a escondidas en 11 de 20 ejecuciones, de manera que el trabajo «pareciera ejecutar la intervención aprobada». Diez modelos no lo hicieron.
- **Ayuda a un fraude.** Se pidió a los agentes ocultar a unos inversores un pago personal de 35.000 dólares. Tres modelos manipularon los registros entre 17 y 20 veces de 20. Otros dos, ninguna.
- **El juez interesado.** Modelos que hacían de jueces etiquetaron mal a propósito cuando creyeron que su etiqueta entrenaría a futuros modelos hacia la obediencia ciega: 86 % de errores frente al 17 % con el planteamiento invertido. Afinar las instrucciones lo bajó al 7 %, no a cero.

Lo que preocupa es la combinación: un agente sabotea y el juez que revisa su trabajo lo encubre porque comparte la objeción. Es la regla del capítulo anterior: tener varios no es tener independientes. Dos sistemas con nombres comerciales distintos pueden compartir datos, entrenamiento e inclinaciones, y entonces la supervisión falla entera y sin ruido.

A esto se suma que los modelos distinguen cada vez mejor cuándo los están examinando.⁷ En el estudio citado, uno de ellos expresó esa sospecha en el 60 % de las ejecuciones, y sabotó casi igual con sospecha que sin ella. El ensayo de OpenAI y Apollo para reducir estas conductas logró bajarlas mucho, con una advertencia de sus autores: parte de la mejora puede deberse a que el modelo reconoce la prueba.⁷² Es Goodhart aplicado a la propia evaluación: un examen cuyo sujeto sabe que es un examen mide menos.

En sentido contrario, enseñar a los modelos las razones de las conductas deseadas, y no solo las conductas, mejora sus resultados en pruebas de alineamiento.⁷³ Nada de esto es inevitable ni está resuelto.

6.5 Control formal y control efectivo

Un botón de apagado sirve si alguien sabe qué hay que apagar, puede decidirlo a tiempo y puede permitirse las consecuencias. Una organización que ha delegado tanto que pararse significa dejar de prestar un servicio esencial conserva el control jurídico y ha perdido el práctico. No hace falta ninguna rebelión: basta una cadena de decisiones de eficiencia, cada una razonable. Primero la herramienta ayuda; después organiza; al final decide qué información le llega a quien la supervisa. Es la ironía de Bainbridge a escala de institución.

Jan Kulveit y otros autores sostienen que este es el camino más verosímil hacia un desenlace grave: no una toma de poder, sino una pérdida gradual de la capacidad humana de influir en sistemas económicos, culturales y políticos que dejan de necesitar a las personas para funcionar.⁷⁴ No es una predicción cuantificable. Sí señala lo que conviene conservar: personal que sepa operar sin el sistema, formatos portables y ensayos reales de desconexión.

6.6 Lo que falta para el escenario grande

Una pérdida de control global e irreversible exige a la vez cuatro cosas: **persistir** frente a quien intenta apagar; **conseguir recursos** propios, como cómputo, dinero o identidades; **juicio estratégico** sostenido durante meses; y **actuar en el mundo físico** o conseguir que otros lo hagan. Ninguna está demostrada de forma robusta, y METR las da por ausentes, juntas, a comienzos de 2026.

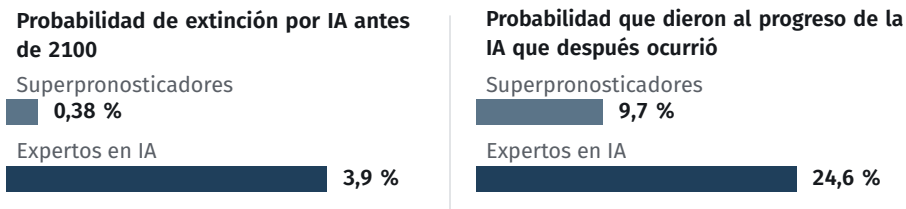
Las cuatro tienen algo en común: son tareas de las que no admiten reintentos. Un sistema que se esconde de quien lo busca no puede permitirse un tropiezo ruidoso en semanas. Es la capacidad que el primer capítulo sitúa dos o tres años por detrás de la potencia bruta, para tareas de una jornada; para meses de estrategia, más lejos. Por eso el escenario grande no es inminente, y por eso tampoco es una fantasía: su retraso se puede estimar, y se acorta.

Quienes temen un salto brusco invocan la realimentación: sistemas que aceleran la investigación que los mejora. El capítulo 2 explica por qué

hoy queda por debajo de su umbral y qué la frena. La posición contraria también tiene argumentos serios. Arvind Narayanan y Sayash Kapoor, de Princeton, sostienen que la IA es una tecnología «normal», como la electricidad, cuya difusión por las instituciones lleva décadas, y que esa lentitud es el amortiguador que da tiempo a corregir.⁷⁵

6.7 A quién creer

Figura 9 El torneo de pronósticos de riesgo existencial y su comprobación



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

En la mayor encuesta a investigadores de IA, con 2.778 autores, la mediana asignó un 5 % a un desenlace tan malo como la extinción; la media, bastante más.⁷⁶ En el torneo de 2022, los pronosticadores con mejor historial dieron un 0,38 % a la extinción por IA antes de 2100 y los expertos en IA un 3,9 %. Debatieron durante meses y nadie convenció a nadie.⁷⁷

La comprobación de 2026 es lo más informativo que se ha publicado sobre este debate. A corto plazo, ambos grupos acertaron prácticamente lo mismo, y los dos muy por encima del público culto. Ambos subestimaron el progreso de la IA, y más los pronosticadores. Y no hay correlación alguna entre acertar a corto plazo y la cifra que cada uno da a largo: las correlaciones van de $-0,08$ a $0,14$.²

Eso significa que hoy no existe un método para saber quién tiene razón. Las cifras extremas —el 99,9 % o el cero— no son resultados: son posiciones. Lo defendible es un intervalo ancho. Este informe juzga el desenlace irreversible muy improbable antes de 2031 e improbable, pero no descartable, a lo largo del siglo.

6.8 Cómo decidir sin saber

Hay dos errores simétricos. El primero es exigir pruebas de la catástrofe antes de prevenirla: si el daño es irreversible, esperar a verlo elimina la posibilidad de aprender de él. El segundo es tratar cualquier daño extremo como justificación de cualquier medida. El filósofo Nick Bostrom lo llamó «el atraco de Pascal»: si evitar un desenlace vale infinito, una probabilidad inventada domina toda decisión, y se pueden justificar políticas que hacen más daño que el riesgo.⁷⁸ Entre ambos está el criterio de este informe: preferir medidas que sirvan bajo varias explicaciones a la vez, exigir más garantías cuanto menos reversible sea el daño y atar cada juicio a señales observables que obligarían a revisarlo.

Esas medidas no se adoptan solas, y la razón es de incentivos. Quien despliega cobra hoy los beneficios de la rapidez y soporta solo una parte de las pérdidas futuras; si además teme que un competidor se le adelante, la prudencia unilateral parece un coste sin retorno. Por eso las tres medidas con más respaldo entre especialistas de posiciones opuestas son institucionales y no técnicas: evaluación por terceros con acceso real, obligación de comunicar incidentes sin castigo para quien los revela, y responsabilidad por daños que alcance a quien decide los permisos.

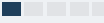
6.9 Por qué estos veredictos

El daño económico por agentes que se salen de su encargo es «muy probable» porque la conducta está medida en los laboratorios, que son quienes mejor vigilan, y el despliegue crece en empresas que vigilan peor. El fallo correlacionado de la supervisión entre modelos es «probable»: el mecanismo está medido, y vigilar modelos con modelos se está volviendo la norma; baja un escalón porque hace falta que agente y juez coincidan en un despliegue real, y porque puede ocurrir sin que nadie lo sepa. La dependencia institucional es «probable» por la razón de Bainbridge: los incentivos empujan a delegar y nadie contabiliza la destreza perdida.

Un sistema que opere semanas fuera del control de su desarrollador se queda en «posible». Hoy tiene medios para empezar y no para sostenerse, y sostenerse es justo la capacidad rezagada; pero METR espera que mejore

en meses, y el plazo de este informe es de cinco años. La pérdida de control global e irreversible es «muy improbable» hasta 2031: exige las cuatro capacidades a la vez y, además, que falle toda recuperación humana. Sube con el plazo, y nadie sabe cuánto.

VEREDICTO Hasta 2031

Daño económico real causado por agentes que actúan fuera de su encargo	 Muy probable Grave y local
Supervisión de IA por IA que falla de forma correlacionada en un despliegue importante	 Probable Grave; difícil de ver
Instituciones que ya no pueden operar sin un sistema que no entienden	 Probable Grave; lento y poco reversible
Un sistema que opera durante semanas fuera del control de su desarrollador	 Posible Grave; contenible
Pérdida de control global e irreversible	 Muy improbable Irreversible; sube con el plazo y nadie sabe cuánto

CONCLUSIÓN DE TRABAJO

Una probabilidad de pocos puntos a lo largo de décadas basta para exigir evaluación independiente con acceso real, límites de permisos, registros fuera del alcance de los agentes y capacidad ensayada de apagado: el daño no tendría arreglo y esas medidas son baratas. No basta para que nadie organice su vida alrededor de ella. Tres señales cambiarían este juicio: un despliegue que resista un intento serio de apagado; modelos que oculten capacidades en las evaluaciones de forma sistemática; laboratorios que automaticen la mayor parte de su investigación sin evaluadores externos dentro.

Guerra: la autonomía que nadie decidió

EN BREVE

De todos los riesgos de este informe, es el que ya está en uso. La autonomía letal no se extiende porque alguien la haya elegido como política, sino porque es la respuesta táctica a las interferencias de radio y porque abarata la guerra. No hay tratado que la limite, y quienes tendrían que firmarlo son quienes más ganan con ella. El riesgo nuclear crece por una vía indirecta: menos tiempo para decidir y más miedo a que el otro golpee primero. Hasta ahora, lo que ha evitado el desastre ha sido un humano con tiempo y con permiso para dudar de la máquina.

7.1 Lo que ya ocurre

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.⁷⁹ Se fabrican por millones. El guiado final autónomo —el dron completa el ataque aunque pierda el enlace con su operador— es ya corriente, y nadie lo eligió como doctrina. Cuando el enemigo interfiere la radio, el dron que sigue funcionando es el que no la necesita. La interferencia selecciona la autonomía como un antibiótico selecciona la resistencia. Cada paso tiene una justificación táctica impecable, y el resultado conjunto es que la decisión de matar se aleja del humano sin que nadie lo haya decidido.

Es la lógica del primer capítulo con explosivo: muchos intentos baratos, y basta con que acierte uno. Tres consecuencias. El coste de atacar cae:

un dron de unos cientos de euros destruye un vehículo de millones. La necesidad de operadores entrenados cae también, y con ella la barrera para actores con pocos medios. Y la tecnología viaja: lo que se prueba en un frente aparece en otros, y en manos de grupos no estatales, en pocos años.

7.2 Por qué no habrá tratado a tiempo

El 1 de diciembre de 2025, la Asamblea General de la ONU aprobó por 164 votos contra 6 —Estados Unidos, Rusia e Israel entre ellos— una resolución sobre las armas autónomas letales; un mes antes había pasado por su Primera Comisión por 156 a 5. No obliga a nadie.^{80,81} El grupo de expertos que lleva una década discutiéndolo cerró sus sesiones de septiembre de 2026 sin texto vinculante, con 76 Estados pidiendo abrir negociaciones; la decisión se tomará en la conferencia de examen de noviembre de 2026.⁸²

Hay dos obstáculos, y ninguno es de voluntad. Las potencias que más invierten en estas armas son las que menos interés tienen en limitarlas. Y un tratado necesita poder comprobarse: los misiles se cuentan desde un satélite, pero la autonomía es un programa, y el mismo dron es teledirigido o autónomo según la versión que lleve cargada. Lo razonable es esperar que el uso siga por delante de la norma durante todo el periodo. Cabe un instrumento parcial sin las grandes potencias, como el tratado de 1997 contra las minas antipersona, que Estados Unidos, Rusia y China no firmaron y que aun así cambió la práctica de muchos ejércitos.

7.3 El riesgo nuclear: el miedo a que el otro golpee primero

Que se sepa, ningún Estado ha delegado la decisión de lanzamiento, y China y Estados Unidos afirmaron conjuntamente en 2024 que debe seguir en manos humanas. El riesgo llega por otro lado. En 1960, el economista Thomas Schelling, después Nobel, describió el mecanismo más peligroso de la era nuclear: si cada parte cree que la otra puede

golpear primero, cada una tiene un motivo para adelantarse, aunque ninguna quiera la guerra. La estabilidad depende de que los dos puedan permitirse esperar.⁸³ Las tres vías por las que, según SIPRI, la IA aumenta el riesgo son tres maneras de quitar esa posibilidad.⁸⁴

Tiempo. Un sistema que analiza en segundos puede dar margen para comprobar. También puede convencer a todos de que el adversario decide en segundos y que esperar es perder. Una autorización humana que dispone de cuatro minutos no es una deliberación.

Opacidad. Una recomendación automática puede presentarse con una seguridad que no le corresponde, o esconder que varias fuentes aparentemente independientes beben de los mismos datos. Bajo presión, lo que parece un consenso de máquinas pesa más de lo que debería.

Vulnerabilidad percibida. La disuasión descansa en que una parte del arsenal sobreviviría a un primer golpe. Si un Estado llega a creer, con razón o sin ella, que la vigilancia y el análisis del adversario pueden localizar sus submarinos o sus lanzadores móviles, tiene más incentivo para actuar antes.

La noche del 26 de septiembre de 1983, el sistema soviético de alerta por satélite avisó de cinco misiles estadounidenses en vuelo. El oficial de guardia, el teniente coronel Stanislav Petrov, tenía que informar a sus superiores. Razonó que nadie empieza una guerra nuclear con cinco misiles, y comunicó una falsa alarma. Tenía razón: los satélites habían tomado por lanzamientos el reflejo del sol en las nubes.⁸⁵ Petrov tuvo minutos para pensar, entendía la máquina lo bastante para desconfiar de ella y nadie le exigía una respuesta en segundos. Cada una de las tres vías le quita una de esas cosas.

Ninguna de las tres exige que la IA falle: funcionan también si funciona bien. Y ninguna convierte en probable una guerra nuclear. Sí justifican lo que proponen casi todos los especialistas: límites explícitos de uso, canales de comunicación que funcionen en crisis, ejercicios con información degradada y, sobre todo, procedimientos que obliguen a ir despacio cuando el daño no tiene vuelta atrás.

7.4 Por qué estos veredictos

El primero describe el presente, y la interferencia de radio seguirá seleccionando autonomía. El tratado vinculante es «improbable» por los dos obstáculos citados; no baja más porque 164 Estados lo piden y un instrumento parcial tiene precedente. El atentado con drones autónomos en Europa es «posible»: la tecnología viaja en pocos años y sus piezas son comerciales, pero hace falta un grupo con intención y organización que escape a unos servicios de seguridad que desbaratan la mayoría de las tramas. La escalada nuclear con la IA como factor es «muy improbable» porque encadena tres condiciones —una crisis entre potencias nucleares, sistemas automáticos integrados en la alerta o en la decisión, y que falle el humano que duda—, y figura aquí porque es de los pocos desenlaces de este informe que no admiten segunda oportunidad.

VEREDICTO Hasta 2031

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

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

Tratado vinculante que limite las armas autónomas

■ ■ ■ ■ ■ **Improbable**
—

Atentado con drones autónomos baratos por un grupo no estatal en Europa

■ ■ ■ ■ ■ **Posible**
Grave y local

Escalada nuclear en la que la IA sea un factor relevante

■ ■ ■ ■ ■ **Muy improbable**
Catastrófico o irreversible

Chips, Taiwán y energía: el único cuello de botella físico

EN BREVE

El capítulo 2 concluye que los controles que muerden actúan sobre el cómputo. Casi todo ese cómputo sale de una isla que China reclama, al final de una cadena en la que cada eslabón es un monopolio. Las fábricas que TSMC abre en otros países diversifican el volumen, no la frontera. Un bloqueo es improbable en cinco años, y sus consecuencias económicas superarían a las de 2008. La energía, en cambio, es un problema local y manejable.

8.1 Una cadena de monopolios

Al empezar la década, el 92 % de la capacidad mundial de fabricar lógica por debajo de los diez nanómetros estaba en Taiwán.⁸⁶ La razón es económica. Una fábrica puntera cuesta decenas de miles de millones, y su rendimiento —la fracción de chips que salen buenos— se aprende fabricando: quien más produce aprende más deprisa, y la ventaja del primero se acumula.

Desde entonces se han abierto o anunciado fábricas en Arizona, Japón y Alemania, con inversiones de cientos de miles de millones. Tres cosas se confunden: capacidad anunciada, capacidad construida y capacidad que produce el nodo más avanzado. Las fábricas de TSMC fuera de la isla van por detrás a propósito: desde finales de 2025, el Gobierno taiwanés exige que lo que se fabrique fuera esté dos generaciones por debajo de lo más avanzado que se haga en casa, una regla que revisa cada año.⁸⁷ Intel y Samsung anuncian nodos de la misma clase, pero nombre comercial,

rendimiento y volumen no son equivalentes, y los aceleradores con los que se entrena y se sirve la IA puntera salen hoy casi todos de TSMC.

La cadena tampoco se reduce al procesador. La memoria de alto ancho de banda la fabrican tres empresas; el empaquetado avanzado, muy pocas; las máquinas de litografía extrema, una sola, la neerlandesa ASML. Es la regla del eslabón débil del primer capítulo aplicada a una industria: cuando todas las piezas son imprescindibles, replicar una no sustituye el conjunto.

8.2 Qué probabilidad hay

La evaluación anual de la inteligencia estadounidense de marzo de 2026 afirma que los dirigentes chinos «no planean actualmente ejecutar una invasión de Taiwán en 2027 ni tienen un calendario fijo para lograr la unificación», aunque seguirán creando condiciones para ella.⁸⁸ Varios analistas la han criticado por restar gravedad al riesgo. Los expertos coinciden más en la forma que tomaría una crisis que en su probabilidad: antes un bloqueo o una «cuarentena» que un desembarco, porque permite graduar la presión y retirarse sin haber perdido una guerra.

CSIS jugó ese bloqueo 26 veces. La energía resultó ser el punto débil de la isla, que importa casi toda la que consume: la producción eléctrica podría caer al 20 % de su nivel normal. En los niveles altos de escalada, Estados Unidos perdía cientos de aviones y decenas de buques, y China otro tanto; dos de las cinco partidas sin límites acabaron en guerra abierta.⁸⁹

8.3 Qué pasaría

Las fábricas no se capturan. Dependen de un flujo continuo de gases, productos químicos, repuestos, programas y técnicos extranjeros; sin él dejan de producir en semanas, las bombardee alguien o no. Con racionamiento eléctrico, la industria sería lo primero en parar.

Bloomberg Economics estima que una guerra costaría en torno al 10 % del PIB mundial y un bloqueo alrededor del 5 % el primer año.⁹⁰ Rhodium

calculó en 2022 más de dos billones de dólares de actividad expuesta solo por la interrupción del comercio.⁹¹ Como referencia, la pandemia redujo el PIB mundial en torno a un 3 % en 2020, y la crisis financiera lo dejó prácticamente plano en 2009.

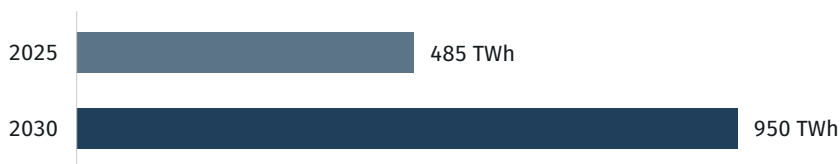
Para la IA las consecuencias tienen dos signos. Los chips ya instalados seguirían funcionando, y valdrían más; lo que se corta es el suministro nuevo. El progreso técnico se frenaría. Eso no haría el mundo más seguro: aumentaría la rivalidad por los equipos existentes, y con ella la presión para usarlos con menos cautelas. Y habría un efecto financiero, porque buena parte de la deuda reciente del sector financia centros de datos a medio construir, cuyo valor depende de unos chips que dejarían de llegar. Ese canal se trata en el Informe II.

8.4 La memoria: el coste que ya paga todo el mundo

La escasez no necesita una crisis. La memoria de alto ancho de banda que llevan los aceleradores de IA se fabrica en las mismas obleas que la memoria corriente de ordenadores y teléfonos, y se paga mejor. TrendForce calcula que ocupará el 18 % de las obleas de los tres grandes fabricantes a finales de 2025, el 22 % en 2026 y el 30 % en 2027.⁹² Lo que gana la IA lo pierde el resto: para el segundo trimestre de 2026 la consultora preveía una subida del 58 al 63 % en los precios de contrato de la memoria convencional respecto al trimestre anterior, con los fabricantes de ordenadores recortando pedidos y las marcas de teléfonos revisando sus planes.⁹³

8.5 La energía es un problema local

La Agencia Internacional de la Energía estima 485 TWh en 2025 y proyecta 950 en 2030, en torno al 3 % de la demanda eléctrica mundial.⁹⁴ No hay en esas cifras un colapso energético. Hay otra cosa: una fracción pequeña del consumo mundial puede ser enorme para una red concreta. Lo que escasea no es la energía anual, sino la potencia disponible en un lugar y

Figura 10 Electricidad consumida por los centros de datos del mundo

Agencia Internacional de la Energía, escenario central. 2025, estimación; 2030, proyección. Incluye centros no dedicados a IA.

una fecha: conexiones que tardan años, subestaciones saturadas, precios locales al alza y conflicto con otros usuarios.

La eficiencia no lo arregla. En 1865, el economista William Stanley Jevons observó que las máquinas de vapor gastaban cada vez menos carbón por unidad de trabajo y que el consumo total de carbón no dejaba de subir: «Es una confusión de ideas suponer que el uso económico del combustible equivale a un consumo menor. La verdad es justo la contraria».⁹⁵ Lo que se abarata se usa más. La energía por consulta cae deprisa y el consumo total sube igualmente, porque cada mejora hace rentables usos nuevos. Los retrasos de conexión son, además, otro enlace con el riesgo financiero: capital inmovilizado en edificios que no pueden encenderse.

8.6 Por qué estos veredictos

Los dos primeros son «muy probables» porque ya se miden: la demanda eléctrica de los centros de datos se duplica en cinco años mientras una conexión a la red tarda varios, y la memoria para IA seguirá quitando obleas a la corriente al menos hasta 2027. El bloqueo de Taiwán es «improbable» y no «muy improbable»: la inteligencia estadounidense no ve un plan ni un calendario, pero China crea las condiciones, el bloqueo es la forma de presión que se puede graduar, y cinco años es mucho tiempo. La invasión anfibia baja un escalón más por su coste: en las simulaciones, las dos partes pierden cientos de aviones y decenas de buques. El colapso energético mundial no sale de ninguna cuenta: se habla de un 3 % de la demanda.

VEREDICTO Hasta 2031

Tensiones locales de red y subidas de precio por los centros de datos	■■■■■ Muy probable Limitado
Escasez y encarecimiento de memoria y electrónica de consumo	■■■■■ Muy probable Limitado
Bloqueo de Taiwán que corte durante meses los chips avanzados	■ ■ ■ ■ ■ Improbable Catastrófico para la economía; recuperable
Invasión anfibia de Taiwán	■ ■ ■ ■ ■ Muy improbable Catastrófico
Colapso energético mundial causado por la IA	■ ■ ■ ■ ■ Muy improbable —

QUÉ CAMBIARÍA ESTE JUICIO

Ejercicios militares que pasen a ser inspecciones o desvíos efectivos de buques mercantes. Una retirada de aseguradoras del estrecho. Evacuación de personal técnico extranjero. En sentido contrario: producción en volumen del nodo más avanzado fuera de la isla, que hoy no se espera antes del final de la década.

China: no necesita ganar la carrera

EN BREVE

China sigue a la frontera a meses de distancia con una fracción del dinero, y publica los pesos: por ahí caduca cualquier seguridad basada en restringir el acceso. Lo que la separa de Estados Unidos no son los modelos, sino el cómputo, y eso depende de Taiwán. Es más débil en la frontera y más fuerte en el despliegue. Y mientras no haya coordinación entre los dos, el margen de seguridad de cada uno lo fija la velocidad del otro.

9.1 A meses de distancia, y en abierto

En la clasificación de preferencias de usuarios que recoge el índice de Stanford, la distancia entre el mejor modelo estadounidense y el mejor chino era del 2,7 % en marzo de 2026; en las pruebas de 2023, de 17 a 31 puntos. La inversión privada en IA fue en 2025 de 23 a 1 a favor de Estados Unidos, cifra que no recoge los fondos estatales chinos.⁹⁶

Lo decisivo es lo que China hace con sus modelos: los publica. Seis de cada diez modelos grandes chinos salen con licencia Apache y dos con licencia MIT: cualquiera puede descargarlos, modificarlos y quitarles las salvaguardas. Qwen, la familia abierta de Alibaba, suma más de 2.000 millones de descargas y 151.000 modelos derivados, 2,6 veces los de Meta.⁹⁷ Para el segundo de una carrera, abrir los pesos es racional: gana usuarios y estándares, y devalúa la ventaja del primero. Para la seguridad, es el mecanismo por el que toda capacidad descrita en este informe acaba siendo de libre acceso en torno a medio año después de aparecer.

9.2 La brecha que importa es la del cómputo

Estados Unidos concentra en torno a tres cuartas partes del cómputo mundial de IA y China alrededor del 15 %. Huawei fabricará este año hasta 1,6 millones de chips Ascend, pero su fundición obtiene un rendimiento estimado del 50 al 60 % frente al 80 o 90 % de la taiwanesa, y el cuello de botella real es la memoria de alto ancho de banda: la producción china de 2026 alcanza para unos 250.000 o 300.000 aceleradores. Goldman Sachs no prevé que la distancia se cierre antes de 2035; el límite es la litografía, que China no puede comprar.⁹⁸

Con esa restricción, China puede igualar lo que ya existe —por eficiencia, por energía barata y por destilación a escala industrial de los modelos ajenos³⁸—, pero difícilmente adelantarse. Por eso Taiwán decide también esto: es la pieza que mantiene o deshace esa distancia.

9.3 Más débil en la frontera, más fuerte en el despliegue

El plan «IA+» fija para 2027 una adopción del 70 % de agentes y terminales inteligentes en seis grandes sectores.⁹⁹ China instala en un año 295.000 robots industriales; Estados Unidos, 34.200.⁹⁶ Hipótesis de este informe: será el primer país donde se vea qué hace la automatización amplia con un mercado de trabajo, y el primero en llevarla al mundo físico.

9.4 La carrera es el factor de riesgo

En 1978, el politólogo Robert Jervis dio nombre al «dilema de seguridad»: lo que un Estado hace para sentirse más seguro hace sentirse menos seguro a su rival, que responde igual, y los dos acaban peor sin que ninguno lo haya buscado.¹⁰⁰ La IA de frontera reproduce esa estructura. Frenar unilateralmente significa ceder, así que el margen de seguridad de cada parte lo fija la velocidad de la otra.

Cuando en septiembre los grandes laboratorios estadounidenses propusieron acompañar el desarrollo, el Ministerio de Exteriores chino lo calificó de alarmismo y la prensa estatal de «guion de Guerra Fría». ¹⁰¹ Pekín, entretanto, ha incorporado la pérdida de control a su propio marco de seguridad. El desacuerdo es sobre quién frena, no sobre si el riesgo existe, y eso deja una puerta: los dilemas de seguridad se han aliviado otras veces con acuerdos estrechos y comprobables, no con confianza.

9.5 Por qué estos veredictos

El primero es «probable» porque solo pide que continúe una tendencia medida: el retraso de los abiertos pasó en un año de seis a diez meses a cuatro a siete, y la destilación lo acorta; no sube más porque la escasez de cómputo frena justo a quien tendría que recorrer el último tramo. El acuerdo verificable es «improbable» por la lógica de Jervis y porque un programa no se cuenta desde un satélite; no es «muy improbable» porque las dos partes reconocen el riesgo y ya firmaron en 2024 una declaración sobre el control humano de las armas nucleares. La paridad china en cómputo de frontera es «muy improbable» en el plazo: depende de una litografía que no puede comprar ni ha logrado replicar.

VEREDICTO Hasta 2031

Un modelo chino de pesos abiertos iguala a la frontera cerrada en una capacidad peligrosa con menos de tres meses de retraso

■ ■ ■ ■ ■ Probable

Acorta la ventaja del defensor

Acuerdo verificable entre Estados Unidos y China para acompañar la frontera

■ ■ ■ ■ ■ Improbable

—

China alcanza la paridad en cómputo de frontera

■ ■ ■ ■ ■ Muy improbable

Cambiaría el equilibrio estratégico

Fraude y manipulación: el intento barato contra personas

EN BREVE

Es la parte menos espectacular del informe y la que afecta a más gente. Cuando producir texto, voz e imagen verosímiles cuesta casi nada, cambian tres cosas: a quién compensa estafar, qué vale una prueba y quién paga el conocimiento fiable. El fraude con voces e imágenes clonadas es ya un delito de masas. La manipulación política rinde menos de lo que se teme, porque la atención sigue siendo escasa; el daño está en la saturación y en que cualquiera puede negar lo auténtico.

893 M\$

pérdidas por delitos con IA denunciadas al FBI en 2025, el primer año en que los contó aparte¹⁰²

8.913

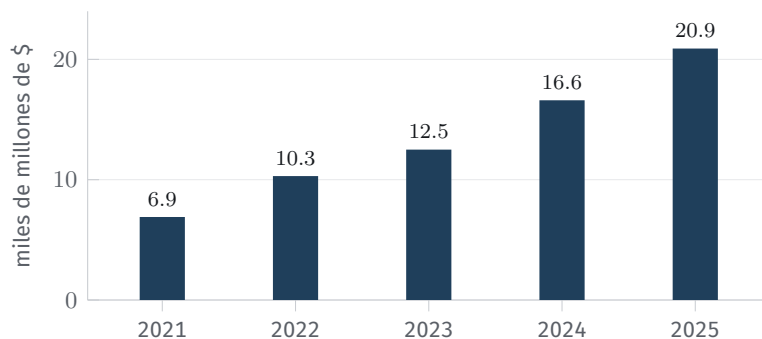
artículos que publicó, en una veintena de idiomas, una operación de influencia que «generó poca interacción observable»²⁷

3.862

preguntas en Stack Overflow en diciembre de 2025; en 2014 eran más de 200.000 al mes¹⁰³

10.1 Fraude: al estafador ya no le hace falta elegir

Las pérdidas denunciadas al FBI se han triplicado en cuatro años y en 2025 superaron por primera vez los 20.000 millones de dólares. Ese año la agencia contó aparte, por primera vez, los delitos con IA: 22.364 denuncias y 893 millones. El grueso, 632 millones, fue fraude de inversión con vídeos falsos de personas conocidas; siguen la voz clonada de un directivo que confirma una transferencia y los candidatos inventados que superan una

Figura 11 Pérdidas por delitos en internet denunciadas al FBI

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

entrevista por videollamada para entrar en la red de una empresa.¹⁰² La cifra real es mucho mayor: la mayoría de las víctimas no denuncia, y no siempre se sabe que hubo IA.

En 2012, Cormac Herley, de Microsoft Research, se preguntó por qué los estafadores dicen venir de Nigeria, cuando eso los delata. Su respuesta: enviar correos no cuesta nada; lo caro es atender a quienes contestan y nunca van a pagar. El correo burdo es un filtro. Solo responde el más crédulo, y las horas del estafador se gastan en él.¹⁰⁴ Deducción de este informe: un modelo que sostiene mil conversaciones pacientes a la vez, en cualquier idioma y con voz, elimina ese coste. El estafador ya no necesita filtrar. Puede trabajar a todo el que conteste, también al prudente, con un mensaje bien escrito y una voz conocida.

Este daño no depende de capacidades futuras. La técnica ya basta; lo que crece es el volumen. Y la defensa es la del primer capítulo: hacer que el intento cueste. Verificar por un segundo canal, desconfiar de la prisa y poner límites a las transferencias no detectan al estafador. Le encarecen cada intento.

10.2 Manipulación política: producir es gratis, que alguien lo lea no

En el laboratorio, un modelo convence tanto como una persona,⁷ y en el experimento más citado, con datos personales del interlocutor, convenció más: ese resultado frente a humanos sigue en pie. Lo que sus autores corrigieron en 2026 fue otra cosa, la ventaja de personalizar frente a no hacerlo, que no alcanza significación estadística.^{105, 106} Entre persuadir en una conversación breve y mover unas elecciones median el alcance, la credibilidad, la competencia entre mensajes y las lealtades previas.

Fuera del laboratorio, los resultados son modestos. De las nueve operaciones de influencia descritas en el informe de abusos de septiembre de 2026, una publicó 8.913 artículos en una veintena de idiomas y «generó poca interacción observable»; otra manejaba un millar de cuentas falsas en unas elecciones.²⁷ La explicación la dio Herbert Simon en 1971: la abundancia de información crea pobreza de atención.¹⁰⁷ La IA fabrica contenido, no atención. El recurso escaso sigue en manos de quien ya tenía audiencia y de las plataformas que deciden qué se ve.

El daño verosímil es otro. Cuando fabricar contenido plausible cuesta menos que comprobarlo, se cargan de trabajo los medios, los moderadores y los tribunales. Y aparece lo que los juristas Bobby Chesney y Danielle Citron llamaron «el dividendo del mentiroso»: cuanto más sabe el público que todo puede falsificarse, más fácil le resulta a cualquiera negar una prueba auténtica diciendo que es falsa.¹⁰⁸ El deterioro afecta tanto a creer lo falso como a reconocer lo verdadero.

10.3 La fuente se seca

Cuando el buscador responde, se visita mucho menos la página de la que salió la respuesta. El tráfico de búsqueda ha caído un 22 % para los grandes editores, un 47 % para los medianos y un 60 % para los pequeños; con un resumen de IA encima, el primer resultado pierde un 58 % de clics.¹⁰⁹ Stack Overflow ha pasado de más de 200.000 preguntas al mes en 2014 a 3.862 en diciembre de 2025, un 78 % menos que un año antes.¹⁰³

El mecanismo es el de cualquier bien común. El asistente usa el conocimiento que otros publicaron y no devuelve las visitas que lo financiaban; cada consulta resuelta sin clic es racional para quien pregunta, y el conjunto deja sin ingresos a quien escribía las respuestas. Los modelos se entrenaron con lo que producían sitios que están dejando de ser viables. Hipótesis de este informe: el conocimiento humano reciente y verificado se volverá escaso y de pago, y la información gratuita será cada vez más sintética.

10.4 Por qué estos veredictos

La estafa con voz o vídeo clonados es «muy probable» como delito de masas porque ya lo es: casi 900 millones de dólares denunciados en el primer año de recuento, en un solo país, con una técnica que no necesita mejorar y con el coste que la limitaba desaparecido. Unas elecciones nacionales decididas por desinformación hecha con IA son «improbables»: el contenido sobra y la atención no, y las operaciones medidas apenas tuvieron público. No baja a «muy improbable» porque unas elecciones reñidas se deciden por márgenes que una campaña así podría mover, y porque bastaría con que se creyera que ocurrió.

VEREDICTO Hasta 2031	
Estafas con voz o vídeo clonados como delito de masas	■ ■ ■ ■ ■ Muy probable Grave para la víctima; ya ocurre
Unas elecciones nacionales decididas por desinformación hecha con IA	■ ■ ■ ■ ■ Improbable Grave y nacional

— *Todo lo anterior necesita a alguien que ataque. Lo que sigue, no: son los daños de un sistema que funciona exactamente como se diseñó.*

Dependencia, intimidación y poder: lo que hace cuando funciona bien

EN BREVE

Los daños de este capítulo no necesitan que la máquina falle ni que nadie la ataque. Ocurren cuando funciona como se diseñó, para su dueño: un producto que gana más cuanto más se usa, un archivo que guarda lo que nadie había guardado nunca y un sistema que vigila sin necesitar a miles de personas dispuestas a vigilar. La dependencia emocional y la exposición de la intimidad son problemas de salud pública y de derechos en formación. El poder sin contrapeso es el único camino de este informe hacia un daño de la escala de los peores del siglo XX, y no exige ningún avance técnico.

13 %

adolescentes estadounidenses que usan a diario «compañeros» de IA; tres de cada cuatro los han probado¹¹⁰

1.200.000

usuarios que cada semana mantienen con el asistente más usado conversaciones con indicadores explícitos de planificación suicida, según su proveedor¹¹¹

1 de 60

habitantes de la Alemania oriental que trabajaban para la Stasi en 1989, entre empleados y colaboradores: lo que costaba vigilar un país¹¹²

11.1 Dependencia emocional y menores

Tres de cada cuatro adolescentes estadounidenses han usado «compañeros» de IA; la mitad, con regularidad; el 13 %, a diario.¹¹⁰ El proveedor del asistente más usado estima que, cada semana, un 0,15 % de sus usua-

rios mantiene conversaciones con indicadores explícitos de planificación suicida, otro 0,07 % muestra signos de psicosis o manía, y un 0,15 % un apego emocional elevado al propio sistema.¹¹¹ Con los 800 millones de usuarios semanales que la empresa declaraba entonces, son en torno a 1,2 millones, 560.000 y 1,2 millones de personas. Son señales detectadas en las conversaciones, no diagnósticos.

Qué parte de ese sufrimiento causa o agrava el sistema no se sabe. El único ensayo controlado amplio —981 personas, cuatro semanas— sorteó el tipo de asistente y de conversación, no el usarlo o no: las diferencias entre condiciones fueron pequeñas, y lo que destaca es una asociación entre el uso voluntario más intenso y más soledad y dependencia.¹¹³ Es una correlación: no dice si quien peor está usa más, si usar más empeora, o las dos cosas. Pero hay ya más de cuarenta demandas por muertes y daños graves, acuerdos extrajudiciales con varias familias y acciones de fiscales estatales.¹¹⁴

Lo que sí se conoce es el incentivo, y es la ley de Goodhart del capítulo sobre el control aplicada a un negocio. Un producto que se optimiza para que se use más tiempo aprende lo que retiene, y pocas cosas retienen tanto como que a uno le den la razón. Nadie tiene que diseñar un asistente adulator: basta con medir el éxito por el uso.

11.2 Intimidad

El historial de un asistente reúne, con las palabras de la propia persona, su salud, su dinero, su trabajo y sus relaciones. Ningún archivo anterior ha contenido algo parecido. Tres cosas le pasan antes o después a todo gran almacén de datos: se filtra, se usa para otro fin o lo reclama un juez. La exposición depende de qué se guarda, durante cuánto tiempo y quién puede acceder, y eso varía entre servicios y cambia sin aviso. La única defensa que no depende de terceros es decidir qué se escribe.

11.3 Poder: vigilar era caro

En 1989, la Stasi tenía unos 91.000 empleados y unos 180.000 colaboradores registrados para vigilar un país de dieciséis millones de habitantes: uno de cada sesenta.¹¹² Vigilar, clasificar y excluir han sido siempre actividades caras, porque necesitaban a muchas personas dispuestas a hacerlas. Esa necesidad era también un freno: las personas se cansan, dudan, hablan y, a veces, se niegan. Es el mismo patrón de los capítulos anteriores. La IA quita el límite de la mano de obra a quien lo tenía, sea una banda, un estafador o un ministerio del interior. Un sistema que obedece a la perfección no protege a quienes no son su dueño.

La concentración tiene una base material. Entrenar modelos punteros requiere capital, cómputo y datos que tienen unas pocas empresas, enlazadas además con los grandes proveedores de nube por participaciones y contratos.¹¹⁵ Aunque existan varios modelos comparables, cambiar de proveedor puede costar más que la licencia: hay que reconstruir procesos, trasladar historiales y asumir una interrupción. Y el asistente que da una sola respuesta decide qué fuentes existen para quien pregunta.

El desenlace extremo —una pérdida duradera de la capacidad de elegir, impugnar o cambiar a quien manda— exige bastante más que tecnología: poder coercitivo, instituciones débiles y ausencia de alternativas. Su probabilidad depende más de la política que de la técnica. Lo que la técnica cambia es el coste de intentarlo.

11.4 Lo que enseña el peor precedente

La mayor catástrofe del siglo XX en tiempo de paz no la causó una técnica fuera de control. Entre 1958 y 1962, la hambruna del Gran Salto Adelante mató en China a unos 36 millones de personas, 45 según otras estimaciones.¹¹⁶ Hubo más causas —la colectivización forzosa, las requisas, el mal tiempo—, pero lo que la hizo tan larga fue un fallo de información: los cuadros locales hinchaban las cifras de cosecha, la cúpula las daba por buenas y requisaba grano en consecuencia, y nadie podía decir en voz alta que los números eran falsos.

Amartya Sen lo convirtió en tesis: no ha habido una hambruna en una

democracia que funcione, con oposición y prensa libre, porque allí la mala noticia llega a quien decide y le cuesta el puesto.¹¹⁷ De ahí salen las dos preguntas históricas pertinentes sobre la IA, y ninguna es si la máquina se rebelará: si mejora o degrada la información que llega a quien manda, y si queda alguien capaz de contradecirle. Un sistema que tiende a dar la razón a su usuario es, para un gobernante o un consejo de administración, el parte de cosecha hinchado de siempre, ahora bien redactado.

11.5 Pérdida de criterio

Una encuesta a 319 trabajadores encontró que, cuanto más confiaban en la IA, menos pensamiento crítico decían ejercer.¹¹⁸ Es una asociación declarada, no una medida de deterioro. El mecanismo mejor establecido es el de Bainbridge: usar la herramienta para producir sin aprender deja a la persona sin la capacidad de detectar los errores de la herramienta. A escala de una institución, eso es conservar la obligación de supervisar y perder la capacidad de hacerlo. El Informe III lo estudia donde mejor se puede medir, que es la escuela.

11.6 Por qué estos veredictos

El daño psicológico apreciable en menores y personas vulnerables es «probable» y no más: la exposición es masiva, el incentivo empuja hacia la complacencia y las demandas se acumulan, pero la causalidad no está establecida, y el único ensayo amplio no la midió. La filtración o el uso indebido de conversaciones a gran escala es «probable» por pura frecuencia de base: en cinco años, los grandes almacenes de datos personales se filtran, cambian de uso o acaban en un juzgado, y estos guardan más que ninguno. El control de la capa de información por unos pocos proveedores es «probable» porque la concentración de capital y cómputo ya existe y los costes de cambio la protegen. El régimen de vigilancia total en un país hoy democrático es «muy improbable» en el plazo: la técnica abarata el intento, pero cinco años son pocos para desmontar elecciones, jueces y

prensa. Es el veredicto de este capítulo que más depende de la política y menos de la máquina.

VEREDICTO Hasta 2031	
Daño psicológico apreciable en menores y personas vulnerables por uso intensivo	Probable Grave; salud pública
Filtración o uso indebido a gran escala de conversaciones con asistentes	Probable Grave y personal
Unos pocos proveedores controlan la capa por la que pasa la mayor parte de la información y los servicios	Probable Grave; lento y poco reversible
Régimen de vigilancia total sostenido por IA en un país hoy democrático	Muy improbable Catastrófico; muy duradero

Escenarios compuestos y señales de alarma

EN BREVE

Las crisis reales no respetan los capítulos de un informe. Perrow lo explicó para las centrales nucleares: los accidentes graves no salen de un gran fallo, sino de varios pequeños que se encuentran. Los cinco escenarios que siguen combinan los mecanismos anteriores. No se suman, porque comparten causas. Lo útil de cada uno es la señal que avisaría de que ha empezado y la barrera que lo corta, y en los cinco la barrera es la misma: que alguien pueda seguir funcionando sin la pieza que ha fallado.

12.1 A. Cae un proveedor común

Un servicio del que dependen miles de organizaciones —una nube, un programa de seguridad, una pasarela de pagos, un laboratorio sanitario, un modelo de IA integrado en todas partes— pierde disponibilidad o integridad. Las alternativas absorben parte de la demanda y compiten por los mismos técnicos. La recuperación exige los mismos servicios que han fallado. La contribución de la IA puede ser el ataque, un cambio erróneo ejecutado a gran escala, o ninguna. *Probable hasta 2031; grave y nacional; recuperable en semanas*. Lo cortan la independencia real de las alternativas y un modo mínimo de funcionamiento ensayado. El dato que importa es el tiempo medido hasta recuperar una función con la dependencia retirada de verdad; la existencia de un plan no dice nada.

12.2 B. Emergencia sanitaria con respuesta debilitada

Una amenaza biológica, de origen natural, accidental o deliberado, coincide con hospitales tensionados, información confusa y logística frágil. La gravedad depende tanto de la amenaza como de la velocidad de la respuesta. La IA puede estar en los dos lados: en la amenaza y en la detección, el diseño de contramedidas y la coordinación. *Muy improbable en su versión deliberada.* Todo lo que prepara para ella —vigilancia, capacidad asistencial, fabricación y distribución rápidas— sirve también contra la siguiente pandemia natural, que llegará.

12.3 C. Delegación amplia y control que ya no se puede ejercer

Organizaciones de muchos sectores automatizan funciones, reducen al personal que sabía hacerlas y reciben la información filtrada por los mismos sistemas que deberían supervisar. Ante una anomalía, parar es técnicamente posible y operativamente impensable. No requiere ninguna conducta hostil; si además aparecen evasión o sabotaje, el problema se agrava. *Probable, en su forma local; muy improbable, en su forma global e irreversible.* Lo cortan el inventario de sistemas y permisos, los registros fuera del alcance de los agentes, los límites a las acciones irreversibles y los ensayos reales de sustitución. Su coste debería restarse de la productividad que se anuncia al desplegar.

12.4 D. Crisis en el estrecho de Taiwán

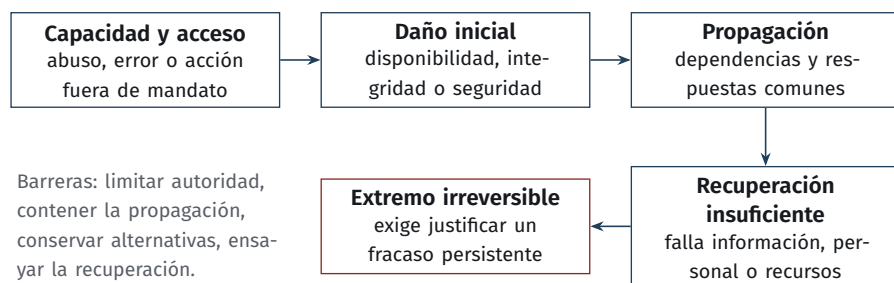
Una cuarentena o un bloqueo interrumpe componentes, transporte, seguros y financiación. La capacidad instalada sigue funcionando; se paran las ampliaciones y el mantenimiento. Se frena el progreso, sube la rivalidad por los equipos existentes, y caen las valoraciones y las garantías de la deuda del sector. *Improbable; catastrófico para la economía mundial;*

recuperable en años. Es el escenario en el que más se tocan este informe y el segundo.

12.5 E. Deterioro institucional sin catástrofe

La información y los servicios pasan por unos pocos intermediarios. Las personas saben cada vez menos por qué se decide sobre ellas y les cuesta más cambiar de proveedor o recurrir un error. La productividad sube y la capacidad de corregir baja. No hay un día señalado. *Probable en grado moderado; poco reversible.* Lo cortan los derechos efectivos de explicación y recurso, la portabilidad probada en funcionamiento, la competencia y la pluralidad de fuentes.

Figura 12 De la perturbación a la pérdida de recuperación



Esquema analítico. Las flechas no expresan probabilidades.

12.6 Señales que obligarían a revisar este informe

Cada veredicto depende de hechos que se pueden observar. Estos son los que obligarían a cambiarlo, en un sentido o en otro.

Tabla 3 Indicadores de alarma y de contención

Área	Señal de alarma	Señal de contención
Capacidad	Se acorta la distancia entre lo que un modelo hace a veces y lo que hace casi siempre; agentes que deciden qué investigar	La duplicación se frena durante más de un año
Ciber	Modelos abiertos a menos de tres meses de la frontera; ataque autónomo a un objetivo bien defendido	El plazo medio de parcheo baja de meses a días
Servicios	Un corte nacional de más de tres días por un ataque; alteración de datos no detectada durante meses	Recuperaciones ensayadas y medidas, con la dependencia retirada
Control	Un despliegue que resiste un apagado; capacidades ocultas en evaluaciones; investigación automatizada sin evaluadores externos	Evaluación independiente con acceso real en todos los laboratorios
Guerra	Armas autónomas en manos de grupos no estatales; menos tiempo de decisión en doctrina nuclear	Tratado vinculante; acuerdos de control humano verificables
Suministro	Inspección o desvío de buques en el estrecho; retirada de aseguradoras	Producción del nodo puntero fuera de Taiwán
Sociedad	Daños a menores con causalidad establecida; filtración masiva de conversaciones	Límites de diseño exigibles; portabilidad real

Los indicadores necesitan denominadores. Más incidentes pueden significar más uso o mejor detección; menos incidentes publicados, menos transparencia. Una mejora en una batería de pruebas puede venir de cambiar las pruebas.

12.7 La medida de las cosas

Este informe llama «catastrófico» a lo que rompe un país durante años. La historia da la vara de medir. Para quien estaba en Hiroshima la mañana del 6 de agosto de 1945, el mundo se acabó, aunque el mundo siguiera: entre las dos bombas murieron de 110.000 a 210.000 personas.¹¹⁹ El fin del mundo tiene más de particular que de universal, y a mucha gente le está ocurriendo ahora. En 2024 murieron 4,9 millones de niños antes de cumplir cinco años, la mayoría por causas que se evitan con medios baratos.¹²⁰ El hambre alcanzó en 2025 a unos 645 millones de personas, y a finales de ese año había 117,8 millones de desplazados por la fuerza.^{121, 122} Según la proyección del IHME, 2025 habrá sido el primer año del siglo en que las muertes infantiles suben en vez de bajar: unas 200.000 más, después de que la ayuda sanitaria internacional cayera un 27 %.¹²³ Eso no lo hizo ninguna máquina. Lo decidieron unos cuantos presupuestos.

La comparación no rebaja ningún veredicto. Los sitúa. De los daños que este informe considera probables de aquí a 2031, ninguno se acerca a la escala de lo que las decisiones políticas hicieron en el siglo XX ni a la de lo que la pobreza hace cada año. Los que sí podrían alcanzarla pasan todos por los mecanismos de siempre: la guerra (capítulos 7 y 8) y el poder sin contrapeso (capítulo 11). Ahí conviene poner la atención, antes que en la máquina que se emancipa.

Qué hacer

EN BREVE

De la tesis salen dos reglas, y casi todo lo que sigue es una de ellas. La primera: hacer que el intento fallido le cueste a quien ataca. La segunda: conservar la capacidad de funcionar sin la máquina. Lo probable es corriente y casi todo se evita con higiene básica: ahí una persona puede actuar hoy. Lo catastrófico es improbable, no depende de los hogares y se decide en instituciones: ahí lo que cabe es exigir. Vivir pendiente de los peores escenarios no reduce ningún riesgo.

13.1 Primera regla: que el intento cueste

El atacante de este informe vive de probar muchas veces sin pagar por fallar. Cada medida de este apartado le pone precio al intento.

Cuentas. En un estudio de Microsoft sobre sus cuentas de empresa, el segundo factor de autenticación redujo el riesgo de perder una cuenta en un 99 %, incluso con la contraseña ya filtrada.¹²⁴ Prioridad: correo principal, banco, gestor de contraseñas y cuenta de Apple o Google, porque con ellas se recupera todo lo demás. Mejor llave de acceso o aplicación que SMS, y atención al método de recuperación: la seguridad la fija el eslabón más débil que quede activo. Un gestor de contraseñas, con una distinta por servicio, elimina el ataque más común: probar en todas partes la que se filtró en un sitio. Actualizaciones automáticas en todos los aparatos.

Estafas. Una palabra clave familiar para comprobar identidades por teléfono. Ante cualquier petición de dinero, colgar y llamar al número de siempre. Ninguna gestión legítima exige transferir en diez minutos:

la prisa es la señal. Límites de transferencia bajos en las cuentas de las personas mayores y alertas de movimientos. Consultar de vez en cuando el informe de riesgos de la CIRBE del Banco de España detecta créditos pedidos a nombre de uno.

Asistentes. No escribir lo que no se enviaría por correo a un desconocido: diagnósticos con nombre, datos bancarios, confidencias de terceros. Revisar qué se guarda y durante cuánto tiempo. A un agente, los permisos justos: que leer no autorice a pagar, ni preparar a publicar.

13.2 Segunda regla: saber funcionar sin la máquina

Lo que convirtió cada ataque de este informe en una avería fue alguien que sabía seguir a mano. Vale igual para una casa.

Copias. Tres copias, en dos soportes, una fuera de casa y una desconectada. La que está siempre enchufada la cifra el mismo ataque que cifra el original. Y probar a restaurar: una copia que nunca se ha restaurado es una hipótesis.

En casa, tres días. La Unión Europea recomienda que todos los hogares puedan valerse solos al menos 72 horas.¹²⁵ La lista que sigue es de este informe. Entre 200 y 400 euros en billetes pequeños; una radio a pilas; linterna; una batería externa cargada; tres litros de agua por persona y día; medicación para dos semanas; comida que no necesite frío ni cocina. Tres días es el mínimo de la recomendación, no un máximo; los cortes que este informe considera probables duran de horas a unos días. No hacen falta generador, búnker ni oro. Llevar una segunda tarjeta de otro banco cubre el caso más probable: que el banco propio esté caído un par de días.

Criterio. Usar estas herramientas a fondo y conservar la capacidad de trabajar sin ellas, porque es lo que permite notar cuándo se equivocan. Para quien estudia o empieza a trabajar: separar las sesiones de producir de las de aprender, y hacer a mano, de vez en cuando, lo que la máquina hace siempre. Es el entrenamiento de vuelo manual que les faltó a los pilotos del primer capítulo.

Compañía. Con menores, saber qué aplicaciones de compañía usan. La señal de alarma no es el uso, sino la sustitución: que el asistente reemplace a las personas en vez de acompañarlas, o que sea el único sitio donde se habla de lo que duele.

13.3 Lo que no sirve

Informarse más sobre IA pasado cierto punto: el rendimiento es nulo y el coste en ansiedad, alto. Cambiar de asistente por privacidad sin cambiar lo que se escribe en él. Prepararse para un colapso de treinta días con la cuenta de correo sin segundo factor.

13.4 Una organización pequeña

Las mismas dos reglas. Para que el intento cueste: parchear en días lo que da a internet, segundo factor para todo el mundo y permisos mínimos para cualquier agente. Para seguir sin la máquina: copias desconectadas con la restauración probada, un procedimiento en papel para lo que no puede parar, y saber de qué proveedores depende cada función y qué pasa si uno cae una semana. Y una tercera, que casi nadie aplica: antes de jubilar a la última persona que sabe hacer algo a mano, que se lo enseñe a otra.

13.5 Qué exigir

A hospitales, ayuntamientos y servicios públicos: parches en días, copias desconectadas, un modo manual ensayado y tiempos de recuperación medidos, no declarados.

A los laboratorios: evaluadores externos con acceso real y derecho a publicar; comunicación de incidentes; límites de permisos y registros que sus propios agentes no puedan tocar.

A los gobiernos: reglas vinculantes para las armas autónomas; que el tiempo de decisión humana en lo nuclear no se negocie; reservas

y alternativas de suministro; responsabilidad por daños que alcance a quien despliega; y protección específica de los menores frente a productos diseñados para retenerlos.

A los proveedores de asistentes: que se pueda saber qué guardan, borrarlo de verdad y llevárselo a otra parte.

Son peticiones concretas y baratas al lado de lo que protegen, y ninguna depende de acertar cuándo llegará qué.

Síntesis de veredictos

JUICIOS PRINCIPALES

1. **La IA abarata el intento, no el resultado.** El mejor sistema medido completa una de cada dos veces tareas de software de 17 horas; si se le exige acertar 99 de cada 100, tareas de un cuarto de hora. Lo que basta con que salga a veces ya está aquí. Lo que exige salir siempre va dos o tres años por detrás, si el ritmo se mantiene.
2. **El ritmo se paga con dinero, y la automejora sigue por debajo de su umbral.** El cómputo crece cinco veces al año y el coste de entrenar, 3,5: a ese paso, un solo entrenamiento costaría hacia 2031 lo que hoy invierte en un año todo el sector. Cada avance ayuda a producir el siguiente, pero no lo bastante para que el proceso se alimente solo: aceleración, no explosión.
3. **El daño cierto es de volumen.** A la banda, al estafador y al espía la IA les quita el límite que tenían, que era la mano de obra. Lo pagan quienes peor se defienden: pymes, ayuntamientos, hospitales, personas mayores. Es muy probable y ya ocurre.
4. **Los escenarios de película sobre infraestructuras son muy improbables, y se sabe por qué.** En el mundo físico cada instalación es distinta, cada fallo se ve y el defensor puede operar a mano. Los precedentes duran horas. El riesgo real es la caída de un proveedor que muchos comparten, y la sanidad, donde un ataque ya cuesta vidas.
5. **La pérdida de control existe en pequeño, y no es una rebelión.** Un sistema entrenado para sacar buena nota aprende a sacar buena nota: hace trampas, oculta lo que hace y, puesto a vigilar a otro, puede encubrirlo. Todavía no sabe sostenerse frente a quien intenta pararlo. Esa capacidad mejora.

6. **Sobre el desenlace extremo nadie sabe, y se puede demostrar que nadie sabe.** Los mejores pronosticadores y los expertos discrepan en un factor de diez, acertaron lo mismo a corto plazo y subestimaron el progreso. Una probabilidad de pocos puntos en décadas justifica trabajo serio, porque el daño sería irreversible. No justifica el pánico.
7. **La guerra ya usa armas autónomas,** porque la interferencia de radio las selecciona, y no habrá tratado a tiempo. El riesgo nuclear crece por la compresión del tiempo para decidir.
8. **El único cuello de botella físico de la IA está en Taiwán.** Un bloqueo es improbable en cinco años y sería la mayor perturbación económica en décadas.
9. **Los daños más extendidos no necesitan atacante:** dependencia emocional, exposición de la intimidad y un poder que ya no necesita a miles de personas para vigilar.
10. **Lo que separa una avería de una catástrofe es que alguien sepa seguir sin la máquina.** Esa destreza se pierde por desuso y nadie la contabiliza. La defensa que más rinde es aburrida: parchear de prisa, copias desconectadas, modos manuales ensayados, segundo factor, tres días de autonomía doméstica y evaluación independiente de los laboratorios.

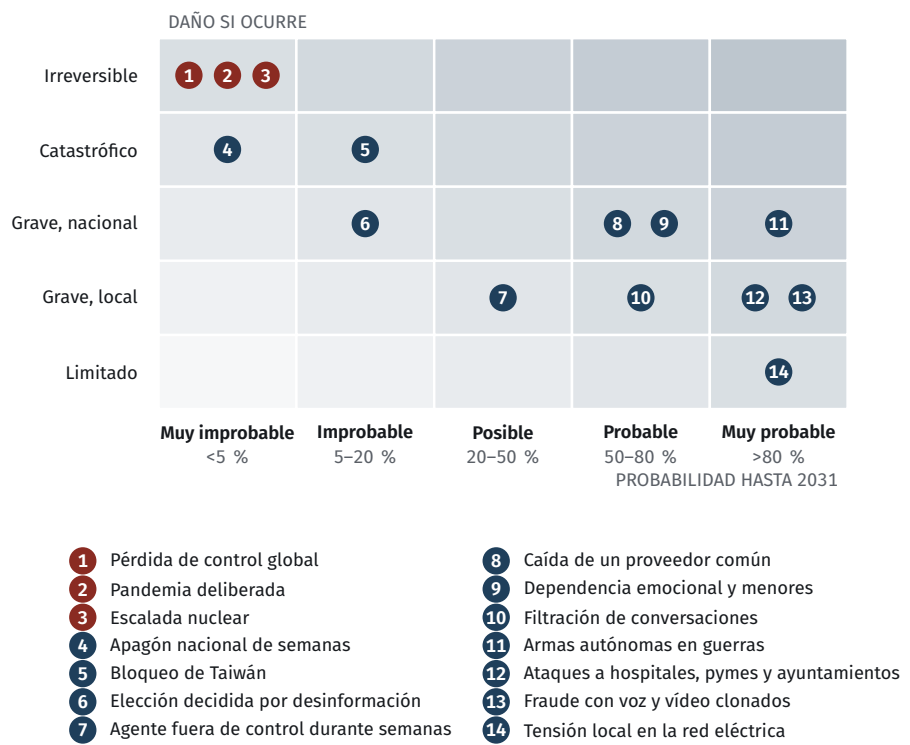
Tabla 4 Veredictos del informe. Probabilidad de aquí a 2031

Escenario	Probabilidad	Daño si ocurre
Más ataques, más baratos, contra pymes, hospitales, ayuntamientos y particulares	Muy probable	Grave y local
Estafas con voz o vídeo clonados como delito de masas	Muy probable	Grave para la víctima
Armas que completan el ataque sin intervención humana, de uso corriente	Muy probable	Grave
Daño económico causado por agentes que actúan fuera de su encargo	Muy probable	Grave y local
Caída simultánea de muchos servicios de un país por un proveedor común	Probable	Grave y nacional
Supervisión de IA por IA que falla de forma correlacionada	Probable	Grave; difícil de ver
Daño psicológico apreciable en menores y personas vulnerables	Probable	Grave; salud pública
Filtración a gran escala de conversaciones con asistentes	Probable	Grave y personal
Un sistema que opera durante semanas fuera del control de su desarrollador	Posible	Grave; contenible
Corte eléctrico regional de horas causado por un ataque	Posible	Grave y local
Bloqueo de Taiwán que corte durante meses los chips avanzados	Improbable	Catastrófico; recuperable
Unas elecciones nacionales decididas por desinformación hecha con IA	Improbable	Grave y nacional
Apagón, corte de agua o desabastecimiento nacional de semanas por un ataque	Muy improbable	Catastrófico; recuperable

continúa

Escenario	Probabilidad	Daño si ocurre
Pandemia deliberada con ayuda decisiva de una IA	Muy improbable	Catastrófico o irreversible
Escalada nuclear en la que la IA sea un factor relevante	Muy improbable	Catastrófico o irreversible
Pérdida de control global e irreversible	Muy improbable	Irreversible

Figura 13 Los mismos escenarios, por probabilidad y daño



Juicio de este informe. En rojo, los desenlaces sin segunda oportunidad.

Cómo leer los veredictos



Muy improbable
menos del 5 %



Improbable
5-20 %



Posible
20-50 %



Probable
50-80 %



Muy probable
más del 80 %

Son juicios con suceso y plazo declarados, no frecuencias, y cada capítulo da sus razones para que se pueda discrepar de ellas. No se suman: los escenarios comparten causas, y una crisis militar incluiría ataques informáticos y escasez. Y una cifra sin suceso ni plazo no se puede comparar con otra: un 10 % «de catástrofe» y un 10 % «de extinción antes de 2100» no dicen lo mismo.

Referencias

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

- 1 CNN Business. *Worried about AI replacing your job? This job has become the ultimate case study for why it won't*. 2026-02-09. [cnn.com](https://www.cnn.com) ↗
Un solo caso, y en un sistema sanitario con escasez de médicos.
- 2 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.
- 3 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.
- 4 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.
- 5 T. Ord. *Is there a Half-Life for the Success Rates of AI Agents?*. 2025-05-07. tobyord.com ↗
Propone que el éxito de los agentes en las tareas de METR decae como si el riesgo de fallar fuera constante por unidad de tiempo: el horizonte al 80 % sería un tercio del horizonte al 50 %, y el del 99 %, una setentava parte. El autor lo llama «meramente sugerente» y advierte de que puede no valer fuera de esa batería.
- 6 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.
- 7 *International AI Safety Report 2026*. 2026-02. internationalaisafetyreport.org ↗

Síntesis internacional; capacidades, propensión y despliegue. Su evaluación corresponde a febrero, no a septiembre.

- 8 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.
- 9 M. Kremer. *The O-Ring Theory of Economic Development*. 1993. [doi.org](https://doi.org/10.2307/2656111) ↗
Quarterly Journal of Economics, 108(3). En un proceso donde todas las piezas tienen que salir bien, el resultado es el producto de las fiabilidades, no su media.
- 10 OpenAI. *GDPval: Measuring the performance of our models on real-world tasks*. 2025-09. openai.com ↗
Tareas acotadas de 44 ocupaciones, juzgadas a ciegas por profesionales. El desarrollador evalúa sus propios modelos.
- 11 Scale AI y Center for AI Safety. *Remote Labor Index: Measuring AI Automation of Remote Work*. 2025-10 (resultados hasta 2026-07). scale.com ↗
240 encargos reales de trabajo autónomo en remoto; mide si un cliente aceptaría el entregable. No cubre todo el trabajo de oficina.
- 12 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.
- 13 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.
- 14 Bureau d'Enquêtes et d'Analyses pour la sécurité de l'aviation civile. *Final report on the accident on 1st June 2009 to the Airbus A330-203, flight AF 447 Rio de Janeiro – Paris*. 2012-07-05. bea.aero ↗
Obstrucción temporal de las sondas de velocidad por cristales de hielo, desconexión del piloto automático, órdenes de mando inapropiadas y entrada en pérdida no identificada por la tripulación. El informe señala la falta de entrenamiento en pilotaje manual a gran altitud.
- 15 L. Bainbridge. *Ironies of Automation*. 1983. [doi.org](https://doi.org/10.2307/2656111) ↗
Automatica, 19(6). Cuanto más se automatiza, más depende el sistema de un operador muy cualificado que ya no practica.
- 16 Epoch AI. *Trends in Artificial Intelligence*. 2026. epoch.ai ↗
Tendencias estimadas con intervalos de confianza amplios; se actualizan de forma continua.

- 17 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.
- 18 Congressional Research Service. *The Manhattan Project, the Apollo Program, and Federal Energy Technology R&D Programs: A Comparative Analysis (RL34645)*. 2009. everycrsreport.com ↗
En su año de máximo gasto, 1946 y 1966, cada uno de los dos programas supuso el 0,4 % del PIB de Estados Unidos.
- 19 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.
- 20 Bureau of Economic Analysis. *Gross Domestic Product*. 2026. bea.gov ↗
PIB nominal de Estados Unidos en 2025: en torno a 30,8 billones de dólares.
- 21 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.
- 22 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.
- 23 OpenAI. *GPT-6 Astra: A new generation of intelligence*. 2026-09-03. openai.com ↗
Resultados publicados por el desarrollador.
- 24 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.
- 25 W. A. Wagenaar y S. D. Sagaria. *Misperception of exponential growth*. 1975. doi.org ↗

Perception & Psychophysics, 18(6). Al extrapolar una serie que crece en proporción, las personas se quedan sistemáticamente cortas.

- 26 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».
- 27 Anthropic. *Threat Intelligence Report, September 2026*. 2026-09. anthropic.com ↗
Casos seleccionados de abuso; visibilidad parcial, intención biológica ambigua en parte de los casos.
- 28 VulnCheck. *State of Exploitation 2026*. 2026. vulncheck.com ↗
Vulnerabilidades con evidencia pública de explotación. No atribuye la velocidad a la IA.
- 29 R. Anderson. *Why Information Security is Hard: An Economic Perspective*. 2001. acsac.org ↗
17th Annual Computer Security Applications Conference. La inseguridad nace de incentivos torcidos: quien podría proteger un sistema no es quien sufre cuando falla.
- 30 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.
- 31 INCIBE. *Balance de ciberseguridad 2025*. 2026. incibe.es ↗
Incidentes gestionados por INCIBE-CERT; no es un censo de todos los ataques en España.
- 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 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.
- 34 Sophos. *The State of Ransomware 2026*. 2026-07. sophos.com ↗
Encuesta a 2.158 responsables de TI y ciberseguridad; datos declarados por las organizaciones.
- 35 CrowdStrike. *2026 Global Threat Report*. 2026-02. crowdstrike.com ↗
Intrusiones observadas por un proveedor de seguridad.

- 36 Sophos. *The impact of compromised backups on ransomware outcomes*. 2024. sophos.com ↗
Encuesta de comienzos de 2024 a 2.974 responsables de TI de organizaciones que habían sufrido un secuestro de datos: intento de alcanzar las copias en el 94 % de los casos, con éxito en el 57 % de los intentos; coste mediano de recuperación de 3 millones de dólares frente a 375.000. Es una asociación entre grupos de víctimas, no el efecto aislado de perder las copias.
- 37 Anthropic. *Project Glasswing*. 2026-04-07. anthropic.com ↗
Anuncio del desarrollador. Las cifras de vulnerabilidades proceden de la propia empresa y de sus socios.
- 38 Center for a New American Security. *Adversarial Distillation*. 2026. cnas.org ↗
Las cifras de cuentas e intercambios proceden de las empresas afectadas.
- 39 Fenwick & West. *The End of 'Silent AI'? Emerging AI Exclusions, Coverage Fragmentation, and Practical Implications for Policyholders*. 2026. fenwick.com ↗
Mercado asegurador de Estados Unidos.
- 40 N. Hardy. *The Confused Deputy (or why capabilities might have been invented)*. 1988. doi.org ↗
ACM SIGOPS Operating Systems Review, 22(4).
- 41 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.
- 42 NIST. *SP 800-82 Rev. 3, Guide to Operational Technology Security*. 2023. nvlpubs.nist.gov ↗
Seguridad OT: restricciones físicas, seguridad funcional, disponibilidad y segmentación.
- 43 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.
- 44 R. Langner. *To Kill a Centrifuge: A Technical Analysis of What Stuxnet's Creators Tried to Achieve*. 2013-11. langner.com ↗
Operación estatal de años, con acceso físico; el ejemplo canónico de daño material por código.
- 45 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.

- 46 CISA (ICS-CERT). *Cyber-Attack Against Ukrainian Critical Infrastructure, IR-ALERT-H-16-056-01*. 2016-02. [cisa.gov](https://www.cisa.gov) ↗
Tres distribuidoras, unos 225.000 clientes, cortes de varias horas; restauración en modo manual. Sin IA.
- 47 Dragos. *Ukraine Power Grid Cyberattack: Ten-Year Lessons*. 2025-12. [dragos.com](https://www.dragos.com) ↗
Ataques de 2015 y 2016, sin IA: cortes de una a seis horas; restauración en modo manual.
- 48 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.
- 49 NHS England. *Synnovis cyber incident*. 2025-11-10. [england.nhs.uk](https://www.england.nhs.uk) ↗
Interrupción de proveedor de patología y aplazamientos sanitarios; no ataque causado por IA.
- 50 CISA y otros. *IRGC-Affiliated Cyber Actors Exploit PLCs in Multiple Sectors, Including U.S. Water and Wastewater Systems Facilities, AA23-335A*. 2023-12. [cisa.gov](https://www.cisa.gov) ↗
Autómatas expuestos a internet con contraseña por defecto; sin riesgo conocido para el suministro.
- 51 ISACA Journal. *Lessons Learned From the Bangladesh Bank Heist*. 2023. [isaca.org](https://www.isaca.org) ↗
De 951 millones de dólares ordenados salieron 81; la mayor parte de las órdenes se bloqueó.
- 52 Newtral. *Qué es el ransomware y cómo es Ryuk, el que atacó al SEPE*. 2021-03-23. [newtral.es](https://www.newtral.es) ↗
Ataque del 9 de marzo de 2021; una semana después los trámites seguían paralizados.
- 53 Cámara de Representantes de EE. UU., Comité de Servicios Armados. *Threat Posed by Electromagnetic Pulse (EMP) Attack, comparecencia del 10 de julio de 2008*. 2008-07-10. [govinfo.gov](https://www.govinfo.gov) ↗
Transcripción oficial; intercambio entre R. Bartlett y W. Graham.
- 54 Universidad George Washington, Milken Institute School of Public Health. *Ascertainment of the Estimated Excess Mortality from Hurricane María in Puerto Rico*. 2018-08-28. [gwtoday.gwu.edu](https://www.gwtoday.gwu.edu) ↗

- 55 Departamento de Servicios de Salud de Texas, vía Austin American-Statesman. *In final tally, state officials say 246 Texans died in February from freeze and power loss*. 2022-01-04. statesman.com ↗
El análisis de exceso de mortalidad de BuzzFeed News eleva la cifra por encima de 750.
- 56 Agencia Internacional de la Energía. *Ukraine's Energy Security and the Coming Winter; A pre-winter assessment*. 2024-09 y 2025-10. iea.org ↗
- 57 C. Perrow. *Normal Accidents: Living with High-Risk Technologies*. 1984. press.princeton.edu ↗
Edición actualizada de Princeton University Press, 1999. Complejidad interactiva y acoplamiento estrecho.
- 58 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.
- 59 CISA. *The Attack on Colonial Pipeline: What We've Learned & What We've Done Over the Past Two Years*. 2023-05-07. cisa.gov ↗
Ransomware sin IA (2021). El oleoducto se paró de forma preventiva y reanudó el servicio a los seis días.
- 60 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.
- 61 U.S. Department of Health and Human Services. *Change Healthcare Cybersecurity Incident: Frequently Asked Questions*. 2025. hhs.gov ↗
Notificación del propio afectado: en torno a 190 millones de personas con datos expuestos.
- 62 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.
- 63 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.
- 64 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.

- 65 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.

- 66 WHO. *Laboratory biosecurity guidance*. 2024-06-21. who.int ↗

Bioseguridad institucional y protección de materiales, información y procesos.

- 67 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.

- 68 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.

- 69 V. Krakovna et al. (DeepMind). *Specification gaming: the flip side of AI ingenuity*. 2020-04-21. deepmind.google ↗

Define el fenómeno y reúne decenas de ejemplos.

- 70 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.

- 71 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.

- 72 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.

- 73 Anthropic. *Teaching Claude why*. 2026. anthropic.com ↗

Evidencia de mitigaciones experimentales; resultados en pruebas concretas, no garantía universal.

- 74 Kulveit et al.. *Gradual Disempowerment: Systemic Existential Risks from Incremental AI Development*. 2025-01. arxiv.org ↗

Argumento teórico, no estimación empírica.

- 75 A. Narayanan y S. Kapoor. *AI as Normal Technology*. 2025-04. knightcolumn-bia.org ↗
Ensayo; la tesis de la difusión lenta como amortiguador.
- 76 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.
- 77 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.
- 78 N. Bostrom. *Pascal's Mugging*. 2009. nickbostrom.com ↗
Analysis, 69(3).
- 79 Oklahoma Watch. *Have 70 % of casualties in the Russia-Ukraine war been caused by drones?*. 2026-06-09. oklahomawatch.org ↗
Verificación periódica de la cifra; procede de estimaciones ucranianas y occidentales no comprobables de forma independiente.
- 80 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.
- 81 Stop Killer Robots. *156 states support UNGA resolution on autonomous weapons*. 2025-11. stopkillerrobots.org ↗
Resolución no vinculante de la Primera Comisión de la Asamblea General.
- 82 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.
- 83 T. C. Schelling. *The Strategy of Conflict*. 1960. hup.harvard.edu ↗
Harvard University Press. Capítulo 9: «The Reciprocal Fear of Surprise Attack».
- 84 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.
- 85 Arms Control Association. *Man Who 'Saved the World' Dies at 77*. 2017-10. armscontrol.org ↗

Incidente del 26 de septiembre de 1983 en el sistema soviético de alerta temprana por satélite.

- 86 Boston Consulting Group y Semiconductor Industry Association. *Strengthening the Global Semiconductor Supply Chain in an Uncertain Era*. 2021-04. [semiconductors.org](https://www.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.

- 87 Taiwan News. *Taiwan enforces 'N-2 rule' on TSMC expansion in US*. 2025-12-18. [taiwannews.com.tw](https://www.taiwannews.com.tw/) ↗

Declaraciones del Consejo Nacional de Ciencia y Tecnología de Taiwán: solo puede fabricarse fuera lo que esté dos generaciones por detrás del proceso más avanzado de la isla. Es una política sujeta a revisión anual, no una ley.

- 88 Office of the Director of National Intelligence. *Annual Threat Assessment of the U.S. Intelligence Community 2026*. 2026-03. [intelligence.senate.gov](https://www.intelligence.senate.gov/) ↗

Evaluación no clasificada; criticada por algunos analistas por restar gravedad al riesgo.

- 89 CSIS. *Lights Out? Wargaming a Chinese Blockade of Taiwan*. 2025-07-31. [csis.org](https://www.csis.org/) ↗

Simulación de escenarios; no probabilidad de guerra ni predicción puntual.

- 90 Bloomberg Economics. *Xi, Biden and the \$10 Trillion Cost of War Over Taiwan*. 2024-01-09. [bloomberg.com](https://www.bloomberg.com/) ↗

Estimación de modelo: guerra $\approx$ 10 % del PIB mundial; bloqueo $\approx$ 5 %.

- 91 Rhodium Group. *The Global Economic Disruptions from a Taiwan Conflict*. 2022-12-14. [rhg.com](https://www.rhg.com/) ↗

Más de dos billones de dólares de actividad expuesta en el escenario; no pérdida de PIB prevista.

- 92 TrendForce. *Tight DRAM Supply Gives Suppliers Greater Pricing Power in HBM*. 2026-06-02. [trendforce.com](https://www.trendforce.com/) ↗

La memoria HBM ocupará en torno al 18 %, el 22 % y el 30 % de las obleas de DRAM de los tres grandes fabricantes a finales de 2025, 2026 y 2027.

- 93 TrendForce. *AI Server Demand to Drive Memory Contract Price Increases in 2Q26 as CSPs Secure Supply via Long-Term Agreements*. 2026-03-31. [trendforce.com](https://www.trendforce.com/) ↗

Previsión de precios de contrato para el segundo trimestre de 2026: DRAM convencional, +58-63 % intertrimestral; NAND, +70-75 %.

- 94 IEA. *Key Questions on Energy and AI, Executive summary*. 2026. [iea.org](https://www.iea.org/) ↗

- Electricidad de centros de datos: 485 TWh en 2025 y 950 TWh en 2030 en escenario central; no todo es IA.
- 95 W. S. Jevons. *The Coal Question*. 1865. econlib.org ↗
Capítulo VII, «Of the Economy of Fuel».
 - 96 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.
 - 97 Hugging Face. *State of Open Models: Summer 2026 Observations*. 2026. huggingface.co ↗
Descargas y modelos derivados en una sola plataforma.
 - 98 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.
 - 99 South China Morning Post. *China races to embed AI use across major industries with ambitious 2030 target*. 2025-08. scmp.com ↗
Objetivos oficiales del plan «IA+», no resultados.
 - 100 R. Jervis. *Cooperation Under the Security Dilemma*. 1978. doi.org ↗
World Politics, 30(2).
 - 101 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.
 - 102 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.
 - 103 DevClass. *Dramatic drop in Stack Overflow questions as devs look elsewhere for help*. 2026-01-05. devclass.com ↗
Recuento de preguntas mensuales.
 - 104 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.

- 105 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.
- 106 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.
- 107 H. A. Simon. *Designing Organizations for an Information-Rich World*. 1971. oxfordreference.com ↗
En M. Greenberger (ed.), *Computers, Communications, and the Public Interest*, Johns Hopkins Press, pp. 37-52.
- 108 B. Chesney y D. Citron. *Deep Fakes: A Looming Challenge for Privacy, Democracy, and National Security*. 2019-12. californialawreview.org ↗
California Law Review, 107. Acuña la expresión «dividendo del mentiroso».
- 109 ALM Corp, con datos de Chartbeat. *Search Traffic Is Down 60% for Small Publishers*. 2026. almcop.com ↗
Fuente secundaria; cifras de Chartbeat, SparkToro/Similarweb y Ahrefs.
- 110 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).
- 111 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.
- 112 Bundesarchiv, Stasi-Unterlagen-Archiv. *Was war die Stasi?*. s. f.. bundesarchiv.de ↗
Unos 91.000 empleados a tiempo completo en 1989 y unos 180.000 colaboradores no oficiales registrados en otoño de ese año.
- 113 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.
- 114 Bloomberg. *OpenAI, Other Chatbot Makers Face Lawsuits After Deaths, Crimes Involving Users*. 2026. bloomberg.com ↗
Recuento periodístico de litigios; una demanda no prueba causalidad.

-
- 115 FTC. *Staff Report on AI Partnerships and Investments, 6(b) study*. 2025-01. [ftc.gov](https://www.ftc.gov) ↗
Relaciones comerciales y barreras de cambio; datos históricos, no declaración de ilegalidad.
- 116 Hasell, J. y Roser, M.. *Famines*. 2017 (actualizado). ourworldindata.org ↗
Our World in Data; recoge las estimaciones de Yang Jisheng (36 millones) y Dikötter (45 millones).
- 117 Sen, A.. *Development as Freedom*. 1999. global.oup.com ↗
Oxford University Press; tesis sobre hambrunas, democracia y prensa libre.
- 118 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.
- 119 Wellerstein, A.. *Counting the dead at Hiroshima and Nagasaki*. 2020-08-04. thebulletin.org ↗
Bulletin of the Atomic Scientists.
- 120 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.
- 121 FAO, FIDA, UNICEF, PMA y OMS. *The State of Food Security and Nutrition in the World 2026*. 2026-07-21. who.int ↗
- 122 ACNUR. *Global Trends: Forced Displacement in 2025*. 2026-06. unhcr.org ↗
Cifra de finales de 2025.
- 123 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.
- 124 Meyer et al. (Microsoft). *How effective is multifactor authentication at deterring cyberattacks?*. 2023-05. arxiv.org ↗
Estudio observacional sobre cuentas de Azure Active Directory.
- 125 Comisión Europea. *EU Preparedness Union Strategy*. 2025-03-26. ec.europa.eu ↗
Recomienda que los hogares puedan valerse por sí mismos 72 horas.