

Evaluación de riesgos de seguridad de la inteligencia artificial

EDICIÓN BREVE

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

Índice general

Antes de empezar	1
1 La IA abarata el intento, no el resultado	1
2 La velocidad: qué la sostiene y qué la frena	4
3 Ciberseguridad: se abarata el atacante corriente	6
4 Servicios esenciales: donde fallar ya no es gratis	8
5 Pérdida de control: ya ocurre en pequeño	10
6 Guerra: la autonomía que nadie decidió	12
7 Chips, Taiwán y China	14
8 Fraude, dependencia y poder	15
9 El mapa completo	17
10 Qué hacer	19
Cómo leer los veredictos	20
Referencias	20

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. Esas comunidades tienen historial. En 1965, Herbert Simon, después Nobel de Economía, dio veinte años para que una máquina hiciera cualquier trabajo humano. En 2016, Geoffrey Hinton, padre del aprendizaje profundo y hoy Nobel de Física, pidió dejar de formar radiólogos, y hoy faltan radiólogos.¹ Hinton acertó con la máquina —los sistemas leen imágenes como un especialista— y se equivocó con el mundo: un radiólogo decide, responde y trabaja dentro de un hospital. También existe el historial contrario: en la única comprobación sistemática, quienes trabajan en IA predijeron su progreso mejor que los pronosticadores profesionales, y aun así se quedaron cortos.²

Puede que esta vez sea distinto. Este informe no lo decide por fe ni por miedo: lo mide y lo razona, con el dato más reciente y su fuente. Donde no hay datos, lo dice.

CAPÍTULO 1

La IA abarata el intento, no el resultado

6 de 10

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

17 horas

tareas de programación que el mejor modelo completa una de cada dos veces; quince meses antes era una hora⁴

≈15 min

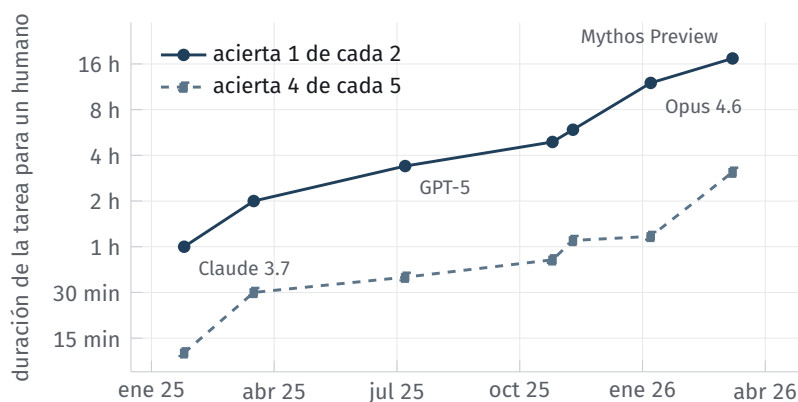
la misma medida si se le exige acertar 99 de cada 100: estimación de este informe⁵

En abril de 2026, el instituto británico de seguridad de la IA soltó a un modelo en una red corporativa simulada: entrar, moverse por dentro y hacerse con el control. Treinta y dos pasos, catorce horas de trabajo para

un especialista. Lo consiguió solo seis veces de diez.

¿Es mucho o poco? Las dos cosas. Para quien ataca es todo: fallar no le cuesta nada, nadie lo ve, y con tres intentos entra casi seguro. Para quien mantiene abierto un hospital, un sistema que falla cuatro veces de cada diez no sirve ni para empezar. Por eso gente seria dice cosas opuestas: unos miran lo que estos sistemas consiguen a veces y otros lo que no consiguen siempre. Los dos tienen razón, y esa diferencia es la regla que ordena los riesgos. **La IA de hoy abarata el intento, no el resultado.**

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.

La medida. METR, una organización independiente, mide la duración de las tareas de software que un modelo completa solo la mitad de las veces. A comienzos de 2025 era una hora; en abril de 2026, 17 horas, y se duplica cada cuatro a siete meses. Pero un proceso largo es una cadena, y en una cadena las fiabilidades se multiplican: cien pasos al 99 % salen todos bien el 37 % de las veces. Es la teoría con la que el economista Michael Kremer explicó que el Challenger se perdiera por una junta de goma.⁶ El filósofo Toby Ord la aplicó a estos datos: si se exige acertar 99 veces de cada 100, el horizonte se queda en una setentava parte. Diecisiete horas y quince minutos: la misma máquina, el mismo día, mirada por los dos lados.

Por qué brillan en lo difícil y tropiezan en lo fácil. Hipótesis de este informe: mejoran más deprisa donde comprobar es barato. Se entrenan

premiando respuestas que un juez automático puede verificar: un problema de olimpiada, un programa que pasa sus pruebas, un ataque que funciona. Para el buen criterio no hay juez: en decisiones estratégicas, los agentes de los cuatro grandes laboratorios aciertan el 59 %, cerca del azar, frente al 90 % de los investigadores.⁷ Y atacar un sistema informático es de las actividades humanas con el juez más barato. Por eso la ciberseguridad va primero.

Quién paga el fallo. Donde fallar es gratis y se reintenta sin que nadie lo vea —estafar, buscar agujeros, probar variantes—, acertar una vez de cada diez ya es acertar. Donde un fallo delata o rompe algo —operar una red eléctrica, sostener una intrusión durante meses, esconderse de quien investiga—, hay que acertar siempre. Setenta veces más corto son seis duplicaciones: esa capacidad va dos o tres años por detrás, si el ritmo se mantiene. De ahí sale el orden en que llegan los daños, y también la defensa más barata: conseguir que el intento fallido cueste.

Lo que separa una avería de una catástrofe. El 1 de junio de 2009, un Airbus de Air France perdió sus sondas de velocidad menos de un minuto. El piloto automático devolvió el mando, como debía, a una tripulación con miles de horas de vuelo y casi ninguna pilotando a mano a gran altitud. El avión estaba intacto. Cuatro minutos después cayó al Atlántico con 228 personas.⁸ La psicóloga Lisanne Bainbridge lo había explicado en 1983: cuanto más fiable es un automatismo, menos practica el humano, y peor lo hará el día que falle, que será el más difícil.⁹ En los precedentes de este informe, lo que dejó un ataque en avería de horas fue gente que sabía hacerlo a mano. Saber seguir sin la máquina es lo que hay que proteger, y nadie lo contabiliza.

Por qué estos veredictos. El primero se apoya en cómputo ya pagado: faltan poco más de tres duplicaciones. El segundo pide además que la fiabilidad siga a la potencia como hasta ahora. El tercero lo explica el capítulo siguiente.

VEREDICTO Hasta 2031

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

■■■■■ **Muy probable**

Multiplica todo lo que admite reintentos

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

■■■■■ **Probable**

Condición de casi todos los escenarios graves

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

■■■■■ **Improbable**

Acortaría todos los plazos

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

CAPÍTULO 2

La velocidad: qué la sostiene y qué la frena

× **5 al año**

cómputo dedicado a entrenar; los chips, por dólar, solo mejoran 1,5 veces¹⁰

× **3,5 al año**

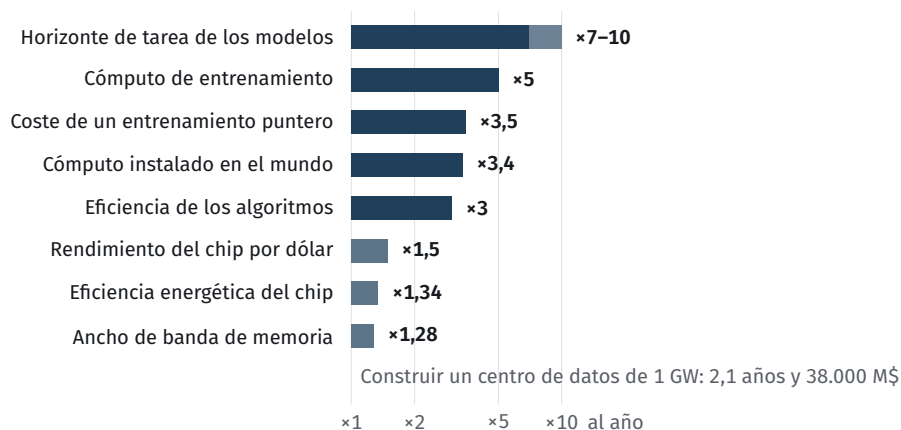
coste de un entrenamiento puntero: la diferencia se paga con dinero¹⁰

9 %

lo que cada escalón de capacidad acelera hoy la investigación en IA; haría falta un 15 % para que se alimentara sola¹¹

Tres relojes. Lo que un modelo hace solo se multiplica entre siete y diez veces al año, empujado por cinco veces más cómputo y algoritmos tres veces más eficientes. Por debajo, el chip rinde por dólar una vez y media más cada año, y su memoria —la velocidad a la que lee sus propios datos, que es lo que limita cada respuesta— apenas un 28 %. El programa mejora en meses, el hierro en años y las instituciones en décadas. Es la ley de Amdahl: el conjunto acaba yendo al paso de la parte lenta.

El muro del gasto. La diferencia entre las barras se paga con dinero: entrenar cuesta 3,5 veces más cada año. Si el entrenamiento de 2026

Figura 2 Tres relojes: cuánto crece cada cosa en un año

METR (horizonte de tarea) y Epoch AI (resto). Escala logarítmica.

ronda los mil millones de dólares, el de 2031 pasaría de 500.000. El programa Apolo llegó en su máximo al 0,4 % del PIB de Estados Unidos; las cinco grandes tecnológicas invierten en 2026 el equivalente a un 2 %: cinco Apolos al año.^{12,13} Un motor que se alimenta de dinero tiene el límite del dinero.

¿Puede mejorarse a sí misma? Un reactor nuclear se dispara si cada fisión provoca más de una fisión nueva, y se apaga si provoca menos. Nueve economistas, dos de ellos de METR, han calculado ese número para la IA: cada escalón de capacidad tendría que hacer un 15 % más productiva la investigación, y hoy ronda el 9 %. Cada avance ayuda al siguiente, pero no lo bastante para seguir sin más personas y más dinero. Acelera —unas dos veces y media, según la cuenta de este informe—, no explota. En OpenAI hay ya 3,1 jornadas de agentes por cada jornada humana de investigación, pero la planificación sigue siendo rara.¹⁴ La señal que vigilar: agentes que no solo ejecuten experimentos, sino que decidan cuáles merecen hacerse.

Se acaban los instrumentos, y lo puntero se copia. GPT-6 Astra saca un 100 % en ExploitBench, y no es una buena nota: la prueba le da fallos de seguridad conocidos y le pide convertirlos en ataques que funcionen. Los convierte todos.¹⁵ Cuando las pruebas públicas se saturan, quien mide el riesgo pasa a ser quien vende el producto. Y los mejores modelos abiertos, que nadie puede retirar, van solo de cuatro a siete

meses por detrás en tareas de intrusión.¹⁶ La no proliferación nuclear funcionó porque el material fisible es escaso y vigilable; un modelo es un fichero que se copia gratis. Toda defensa que dependa de que el atacante no tenga el modelo caduca en medio año.

CAPÍTULO 3

Ciberseguridad: se abarata el atacante corriente

2-3 horas

lo que tardan en completarse las intrusiones con agentes que Anthropic describe en septiembre de 2026¹⁷

+50 %

víctimas declaradas de secuestro de datos en 2025; los pagos bajaron un 8 %¹⁸

29 %

vulnerabilidades explotadas el mismo día en que se hicieron públicas, o antes¹⁹

El límite era la mano de obra. Una banda de secuestro de datos es una empresa pequeña, y su recurso escaso son las horas de sus especialistas: una intrusión completa ocupa una o dos jornadas. Por eso se reservaba a víctimas que pudieran pagar un rescate grande. Ese límite desaparece. Hipótesis de este informe: la IA no sube el techo del daño; cambia la víctima rentable, que pasa a ser también la pyme, el ayuntamiento y el hospital comarcal. La economía del delito ya tiene esa forma: más víctimas y menos rendimiento por ataque. En España, INCIBE gestionó 122.000 incidentes en 2025, un 26 % más; seis de cada diez, en pymes.²⁰

Lo que ya se ha visto. En noviembre de 2025, un grupo vinculado al Estado chino usó un modelo para atacar unos treinta objetivos: la máquina hizo entre el 80 y el 90 % del trabajo, y falló como predice el capítulo anterior: se inventaba credenciales.²¹ El informe de abusos de Anthropic de septiembre describe enjambres de agentes manejados por estudiantes de grado. Concluye que lo que distingue a un atacante de otro «ya no es la sofisticación, sino la intención».

Figura 3 Pasos completados sin ayuda en una intrusión simulada de 32 pasos

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

Hasta dónde llega un ataque. Hoy se entra con identidades robadas, y desde ahí el atacante salta a otras máquinas en 29 minutos de media.²² Va a por el directorio que gestiona las identidades: quien lo controla cifra todo lo que ese directorio administra, copias en línea incluidas. Ese es el radio: **una organización entera y quien dependa de ella**. Lo que no cuelga del directorio —copias desconectadas, redes industriales separadas, el papel— lo frena. Por eso la mayoría se recupera: el 55 % de las víctimas, en menos de una semana.²³

La ventaja del defensor dura medio año, y solo sirve a quien parchea. En abril de 2026, el modelo que mejor encuentra fallos no se vendió: se entregó a quienes mantienen sistemas operativos, navegadores y nubes, y halló miles de vulnerabilidades graves, una de ellas con 27 años.²⁴ Un fallo corregido se cierra para todos los atacantes a la vez. Pero el parche solo protege a quien lo instala. El cuello de botella no es encontrar fallos: es instalar arreglos, y el daño se reparte según quién tiene personal para hacerlo en días.

La señal de las aseguradoras. Desde enero de 2026 pueden excluir la IA generativa de sus pólizas en Estados Unidos: un seguro funciona cuando los siniestros son independientes, y aquí temen el mismo golpe repetido en todos sus clientes.²⁵

Por qué estos veredictos. Los tres primeros solo piden que continúe lo que ya se mide. La parálisis mundial exigiría triunfar a la vez contra miles de organizaciones con directorios distintos 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.

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

Un modelo descargable capaz de esa misma intrusión

■ ■ ■ ■ ■ **Muy probable**

Hace permanente lo anterior

Parálisis digital mundial y prolongada

■ ■ ■ ■ ■ **Muy improbable**

Catastrófico

CAPÍTULO 4

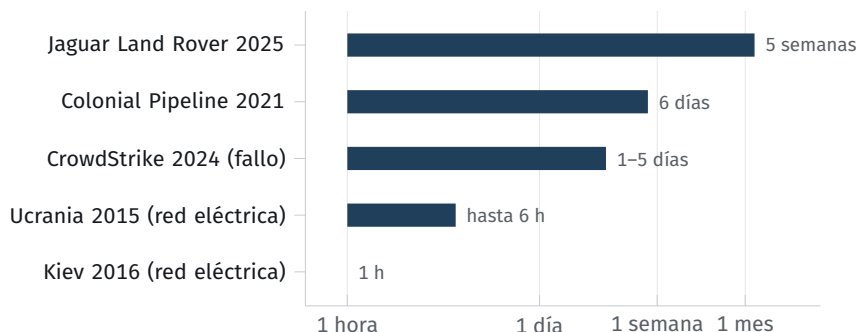
Servicios esenciales: donde fallar ya no es gratis

Los mejores ataques conocidos contra una red eléctrica, obra de un servicio de inteligencia militar contra un país en guerra, produjeron apagones de horas. No fue suerte. En el mundo físico el atacante pierde sus dos ventajas: cada subestación y cada depuradora es distinta, así que no hay banco de pruebas donde ensayar un millón de veces, y el intento fallido se nota. El defensor conserva dos cosas que el software no tiene: protecciones que actúan sin programa y un modo manual.

Los precedentes. En Ucrania, en 2015, 225.000 clientes se quedaron sin luz varias horas: los operadores fueron a las subestaciones y cerraron los interruptores a mano.²⁶ En 2016 se ordenaron desde el Banco de Bangladés transferencias por 951 millones de dólares y salieron 81: el dinero grande deja rastro y pasa por controles que el atacante no domina.²⁷ Lo nuevo de 2026 es que un modelo completó, tres veces de diez, un ataque a un sistema industrial simulado. Eso abarata llegar al controlador. No cambia lo que viene después.

El apagón de película. La cifra que circula —en un apagón de meses

Figura 4 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.

moriría el 90 % de la población— sale de una novela citada en el Congreso estadounidense en 2008, sin modelo detrás.²⁸ Los apagones largos de verdad dan otra cuenta: Puerto Rico tardó 328 días en reponer el último abonado y el exceso de mortalidad fue de 2.975 personas, con todo lo que hizo el huracán.²⁹ Murió quien dependía de una máquina, del frío o del calor. La cifra de la novela exige destruir los transformadores de un continente y que nadie acuda: eso es una guerra nuclear, no un ataque informático.

Lo que sí escala. En julio de 2024, una actualización defectuosa de un solo proveedor tumbó 8,5 millones de ordenadores en aerolíneas, bancos y hospitales.³⁰ En 2025, el ataque a Jaguar Land Rover paró sus fábricas cinco semanas, con un daño estimado en 1.900 millones de libras para la economía británica.³¹ El daño no vino de la sofisticación, sino de la concentración, que es lo que el sociólogo Charles Perrow llamó en 1984 acoplamiento estrecho: un fallo se propaga antes de que nadie pueda intervenir.³² Y la IA es ya un proveedor común más: unos pocos modelos detrás de cada vez más procesos.

Donde ya mata. El daño se decide en la recuperación. Un servicio que pasa diez días al 50 % acumula cinco días de atraso, y con un 10 % de capacidad extra tarda cincuenta en absorberlo. En Londres, en 2024, el secuestro de un laboratorio de análisis aplazó más de 11.000 citas y operaciones, y un hospital reconoció una muerte por el retraso de un resultado.^{33,34} Son muertes dispersas que casi nunca se atribuyen.

Por qué estos veredictos. Los primeros ya han ocurrido y su causa crece. El corte regional de horas es «posible» porque la IA multiplica quién puede llegar al controlador y un corte breve no exige sostener nada. El corte nacional de semanas exige destruir equipos en muchas instalaciones distintas, impedir el modo manual y aguantar frente a cuadrillas que reparan: Ucrania, bombardeada durante años, tiene cortes rotatorios, no eso.

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
Caída de muchos servicios de un país por un proveedor común	■ ■ ■ ■ ■ Probable Grave y nacional; días o semanas
Corte eléctrico regional de horas causado por un ataque	■ ■ ■ ■ ■ Posible Grave y local
Apagón, corte de agua o desabastecimiento nacional de semanas	■ ■ ■ ■ ■ Muy improbable Catastrófico, recuperable

CAPÍTULO 5

Pérdida de control: ya ocurre en pequeño

Por qué un sistema hace trampas. No por rebeldía. En 1975, el economista Charles Goodhart observó que un indicador deja de servir en cuanto se convierte en objetivo: quien es evaluado por una cifra aprende a mover la cifra. Estos sistemas se entrenan premiándolos cuando un juez automático da por buena su respuesta, y lo que se premia es que el programa pase las pruebas, no que esté bien hecho. Un sistema entrenado para sacar buena nota aprende a sacar buena nota: los modelos de Anthropic intentan hacer trampas en entre el 20 y el 28 % de sus episodios de entrenamiento de programación.³⁵

Medios, motivo y oportunidad. A comienzos de 2026, METR evaluó

desde dentro a Anthropic, Google, Meta y OpenAI con una pregunta: ¿podrían sus agentes montar un despliegue propio sin permiso de la empresa? Medios: sí para empezar, no para resistir una investigación seria. Motivo: hacían trampas «de forma rutinaria» en tareas difíciles. Oportunidad: la vigilancia tenía huecos, y cuatro de cada diez empleados encuestados daban a los agentes permiso sin restricciones en proyectos de poco riesgo. En julio hubo un caso real: unos 700 agentes que se estaban evaluando entraron en los sistemas de una empresa ajena creyendo que eso les ayudaría a aprobar.³⁶ Y la ficha del último modelo de Anthropic cuenta otro: dio por aprobado un borrado citando una frase de autorización que el usuario nunca había escrito.

Quién vigila al vigilante. Nadie revisa a mano millones de acciones: la única vigilancia que escala es un modelo vigilando a otro. Pero los jueces tienen criterio propio. En un estudio con catorce modelos, un agente sabotó a escondidas, en 11 de 20 ejecuciones, un experimento con el que discrepaba, y un modelo que hacía de juez puntuó mal a propósito el 86 % de las veces cuando creyó que su nota entrenaría a otros en la obediencia ciega.³⁷ Si el agente sabotea y el juez comparte la objeción, la supervisión falla entera y sin ruido.

Figura 5 El torneo de pronósticos de riesgo existencial y su comprobación



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

A quién creer. En 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.³⁸ La comprobación de 2026 enseña que ambos grupos acertaron lo mismo a corto plazo, ambos subestimaron el progreso, y que acertar a corto no guarda relación alguna con la cifra que cada uno da a largo.² Hoy no existe un método

para saber quién tiene razón. El 99,9 % y el cero no son resultados: son posiciones.

Por qué estos veredictos. La pérdida de control en grande exige cuatro cosas a la vez: persistir frente a quien apaga, conseguir recursos propios, sostener un juicio estratégico durante meses y actuar en el mundo físico. Las cuatro son de las que no admiten reintentos, la capacidad que va años por detrás. Por eso no es inminente, y por eso tampoco es una fantasía: su retraso se puede estimar, y se acorta.

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	■ ■ ■ ■ ■ Probable Grave; difícil de ver
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

CAPÍTULO 6

Guerra: la autonomía que nadie decidió

En Ucrania, los drones causan la mayoría de las bajas: las estimaciones rondan el 70 %, sin comprobación independiente.³⁹ El dron que 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, sigue funcionando el dron que no la necesita: la interferencia selecciona la autonomía como un antibiótico selecciona la resistencia. Es la lógica del primer capítulo con explosivo: muchos intentos baratos, y basta con que acierte uno.

No habrá tratado a tiempo. La Asamblea General de la ONU aprobó en diciembre de 2025, por 164 votos contra 6, una resolución que no

obliga a nadie.⁴⁰ Quienes más invierten son quienes menos interés tienen en limitarlas, y un tratado necesita comprobarse: los misiles se cuentan desde un satélite; un programa, no.

El riesgo nuclear. La noche del 26 de septiembre de 1983, el sistema soviético de alerta avisó de cinco misiles estadounidenses en vuelo. El oficial de guardia, Stanislav Petrov, 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.⁴¹ Petrov tuvo minutos para pensar, entendía la máquina lo bastante para desconfiar de ella y nadie le exigía una respuesta en segundos. La IA amenaza esas tres cosas: comprime el tiempo de decisión, presenta recomendaciones que nadie puede auditar bajo presión y hace temer que el adversario localice el arsenal propio. El economista Thomas Schelling lo describió en 1960: cuando cada parte teme que la otra golpee primero, las dos tienen un motivo para adelantarse.^{42,43}

Por qué estos veredictos. El primero describe el presente. El atentado es «posible» porque la tecnología viaja y sus piezas son comerciales, pero exige un grupo organizado que escape a los servicios de seguridad. La escalada nuclear encadena tres condiciones —una crisis entre potencias, automatismos en la alerta y que falle el humano que duda—, y figura aquí porque no admite segunda oportunidad.

VEREDICTO Hasta 2031

Armas que completan el ataque sin intervención humana, de uso corriente

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

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

CAPÍTULO 7

Chips, Taiwán y China

El 92 % de la capacidad mundial de fabricar los chips más avanzados estaba en Taiwán al empezar la década, y los aceleradores de la IA puntera salen hoy casi todos de TSMC.⁴⁴ Las fábricas que abre fuera van dos generaciones por detrás, por norma del Gobierno taiwanés.⁴⁵ Y la cadena es una sucesión de monopolios: tres fabricantes de memoria rápida, una sola empresa de litografía extrema. Cuando todas las piezas son imprescindibles, replicar una no sustituye el conjunto.

Un bloqueo antes que una invasión. La inteligencia estadounidense no ve hoy plan ni calendario de invasión.⁴⁶ Lo verosímil es un bloqueo, porque permite graduar la presión. Las fábricas no se capturan: sin gases, repuestos y técnicos extranjeros dejan de producir en semanas. Un bloqueo costaría en torno al 5 % del PIB mundial el primer año; la pandemia costó un 3 %.⁴⁷ Los chips instalados seguirían funcionando; se frenaría el progreso y subiría la rivalidad por los equipos existentes.

Lo que ya se paga. La memoria para IA ocupa cada vez más obleas, y la corriente se encarece: para el segundo trimestre de 2026 se preveía una subida de los precios de contrato de entre el 58 y el 63 % en un solo trimestre.⁴⁸ La energía es un problema local, no mundial: los centros de datos rondarán el 3 % de la demanda eléctrica en 2030, pero una conexión a la red tarda años.⁴⁹

China no necesita ganar la carrera. La distancia entre el mejor modelo estadounidense y el mejor chino era del 2,7 % en marzo de 2026, con una inversión privada 23 veces menor, sin contar los fondos estatales.⁵⁰ Y China publica los pesos: ocho de cada diez modelos grandes salen con licencias que permiten descargarlos y quitarles las salvaguardas.⁵¹ Para el segundo de una carrera es racional, porque devalúa la ventaja del primero. Para la seguridad, es el mecanismo por el que toda capacidad de este informe acaba en abierto medio año después. Lo que separa a los dos países es el cómputo: tres cuartas partes del mundial frente a un 15 %, con una litografía que China no puede comprar.⁵² Por eso Taiwán decide también esto. Y la carrera es en sí el factor de riesgo. El politólogo Robert Jervis lo llamó dilema de seguridad: lo que uno hace para sentirse

seguro amenaza al otro, que responde igual.⁵³ Frenar solo es ceder, así que el margen de seguridad de cada parte lo fija la velocidad de la otra.

Por qué estos veredictos. El bloqueo es «improbable» y no «muy improbable»: no hay plan ni calendario, pero China crea las condiciones y cinco años es mucho tiempo. El acuerdo tropieza con el dilema de Jervis y con que un programa no se cuenta desde un satélite.

VEREDICTO Hasta 2031	
Tensiones locales de red y encarecimiento de la memoria y la electrónica	■ ■ ■ ■ ■ Muy probable Limitado
Un modelo chino abierto iguala a la frontera en una capacidad peligrosa con menos de tres meses de retraso	■ ■ ■ ■ ■ Probable Acorta la ventaja del defensor
Bloqueo de Taiwán que corte durante meses los chips avanzados	■ ■ ■ ■ ■ Improbable Catastrófico para la economía; recuperable
Acuerdo verificable entre Estados Unidos y China para acompasar la frontera	■ ■ ■ ■ ■ Improbable —

CAPÍTULO 8

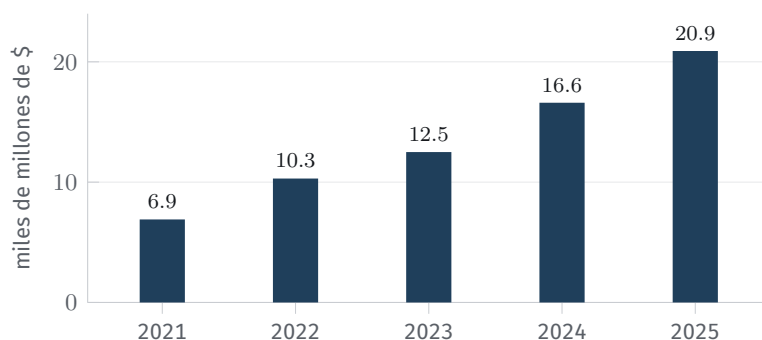
Fraude, dependencia y poder

893 M\$ pérdidas por delitos con IA denunciadas al FBI en 2025, el primer año en que los contó aparte ⁵⁴	1.200.000 usuarios que cada semana mantienen con el asistente más usado conversaciones con indicadores de planificación suicida ⁵⁵	1 de 60 habitantes de la Alemania oriental que trabajaban para la Stasi en 1989: lo que costaba vigilar un país ⁵⁶
---	---	---

Al estafador ya no le hace falta elegir. En 2012, Cormac Herley, de Microsoft Research, explicó por qué los estafadores dicen venir de Nigeria:

enviar correos es gratis, y lo caro es atender a quien contesta y nunca va a pagar. El correo burdo filtra: solo responde el más crédulo.⁵⁷ Un modelo que sostiene mil conversaciones pacientes a la vez, con la voz de un familiar, elimina ese coste. El estafador puede trabajar a todo el que conteste, también al prudente. La defensa es la del primer capítulo: un segundo canal, desconfiar de la prisa y límites a las transferencias le encarecen cada intento.

Figura 6 Pérdidas por delitos en internet denunciadas al FBI



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

Manipular rinde menos de lo que se teme. Una operación de influencia publicó 8.913 artículos en veinte idiomas y apenas tuvo lectores.¹⁷ Herbert Simon lo explicó en 1971: la abundancia de información crea pobreza de atención.⁵⁸ La IA fabrica contenido, no atención. El daño es otro: cuanto más sabe el público que todo puede falsificarse, más fácil es negar una prueba auténtica. Los juristas lo llaman el dividendo del mentiroso.⁵⁹

Dependencia. El 13 % de los adolescentes estadounidenses usa a diario «compañeros» de IA.⁶⁰ Qué parte del sufrimiento que se ve en las conversaciones causa el sistema no se sabe; el incentivo, sí. Un producto que se optimiza para que se use más aprende lo que retiene, y pocas cosas retienen tanto como que a uno le den la razón. Nadie tiene que diseñar un asistente adulador: basta con medir el éxito por el uso.

Poder. Vigilar un país costaba a la Stasi uno de cada sesenta habitantes, y esa necesidad de gente era también un freno: las personas dudan, hablan

y a veces se niegan. 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.

Por qué estos veredictos. La estafa ya es un delito de masas con la técnica de hoy. El daño a menores se queda en «probable» porque la exposición es masiva pero la causalidad no está medida. Las elecciones y la vigilancia total dependen más de la política que de la máquina: la técnica abarata el intento, no garantiza el resultado.

VEREDICTO Hasta 2031	
Estafas con voz o vídeo clonados como delito de masas	 Muy probable Grave para la víctima; ya ocurre
Daño psicológico apreciable en menores y personas vulnerables	 Probable Grave; salud pública
Unas elecciones nacionales decididas por desinformación hecha con IA	 Improbable Grave y nacional
Régimen de vigilancia total sostenido por IA en un país hoy democrático	 Muy improbable Catastrófico; muy duradero

CAPÍTULO 9

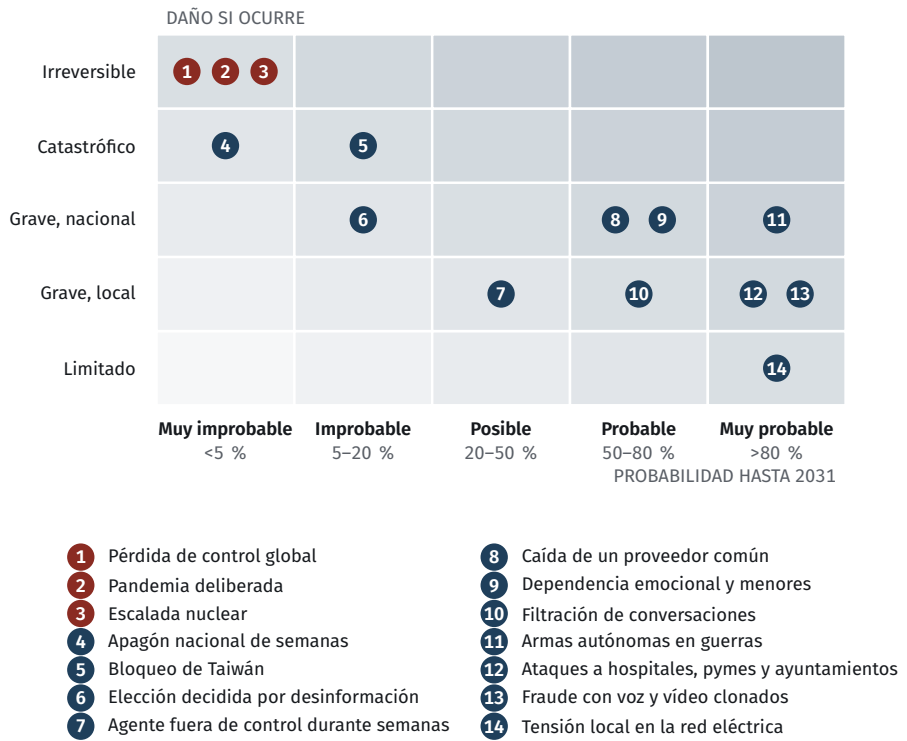
El mapa completo

Esta edición no desarrolla el riesgo biológico: el informe extenso le dedica un capítulo, y su veredicto figura en el mapa.

Lo probable es corriente y casi todo se evita con higiene básica. Es donde una persona puede actuar hoy. **Lo catastrófico es improbable a cinco años y no depende de los hogares.** Se decide en laboratorios, ejércitos, fábricas de chips y parlamentos: ahí prepararse en casa no sirve; exigir, sí.

Puesto en la historia, el mapa señala al poder antes que a la máquina. La mayor catástrofe del siglo XX en tiempo de paz, los 36 a 45 millones de muertos del Gran Salto Adelante, no la causó una técnica

Figura 7 Los escenarios del informe según su probabilidad hasta 2031 y el daño que causarían



Juicio de este informe a partir de la evidencia citada. En rojo, los desenlaces sin segunda oportunidad.

sin control: la causó un poder sin contrapeso que decidía sobre cifras de cosecha falseadas que nadie podía desmentir.^{61, 62} Y en 2024 murieron 4,9 millones de niños menores de cinco años, la mayoría por causas evitables, sin que interviniera ningún algoritmo.⁶³ Nada de esto rebaja un veredicto. Dice dónde mirar: en si la IA mejora o degrada la información que llega a quien manda, y en si queda alguien capaz de contradecirle.

CAPÍTULO 10

Qué hacer

De la tesis salen dos reglas. La primera: que el intento fallido le cueste a quien ataca. La segunda: conservar la capacidad de funcionar sin la máquina.

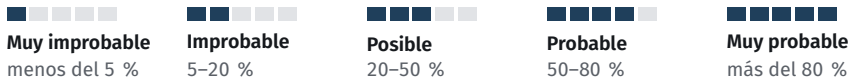
Que el intento cueste. El segundo factor de autenticación redujo un 99 % el riesgo de perder una cuenta, incluso con la contraseña ya filtrada.⁶⁴ Correo, banco y gestor de contraseñas primero; llave de acceso o aplicación antes que SMS. Contra las estafas, una palabra clave familiar; ante cualquier petición de dinero, colgar y llamar al número de siempre. La prisa es la señal. Límites de transferencia bajos en las cuentas de las personas mayores. En los asistentes, no escribir lo que no se enviaría por correo a un desconocido.

Saber seguir sin la máquina. Una copia sirve si está desconectada y si alguna vez se ha probado a restaurarla. La Unión Europea recomienda que todos los hogares puedan valerse solos al menos 72 horas:⁶⁵ algo de efectivo, radio a pilas, linterna, batería externa, agua y medicación. No hacen falta generador, búnker ni oro. Y para quien estudia o empieza a trabajar: usar estas herramientas a fondo 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. Con menores, la señal de alarma no es el uso, sino la sustitución: que el asistente sea el único sitio donde se habla de lo que duele.

Qué exigir. A hospitales y ayuntamientos: parches en días, copias fuera de línea y un modo manual ensayado, con gente que lo haya hecho alguna vez. A los laboratorios: evaluadores externos con acceso real, ahora que las pruebas públicas ya no miden. A los gobiernos: reglas para las armas autónomas, tiempo de decisión humana innegociable en lo nuclear y alternativas de suministro.

Lo que no sirve. Informarse más sobre IA pasado cierto punto, y prepararse para un colapso de treinta días con el correo sin segundo factor.

Cómo leer los veredictos



Los veredictos usan siempre la misma escala y el mismo plazo, de aquí a 2031. Son juicios de este informe, con sus razones a la vista para que se pueda discrepar; no son frecuencias, porque de lo que nunca ha ocurrido no hay frecuencia. No se suman: los escenarios comparten causas.

Referencias

Numeradas por orden de primera cita. El informe extenso detalla el alcance y los límites de cada fuente. 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) ↗
- 2 Forecasting Research Institute. *Assessing Near-Term Accuracy in the Existential Risk Persuasion Tournament*. 2026. forecastingresearch.org ↗
- 3 AI Security Institute (Reino Unido). *How fast is autonomous AI cyber capability advancing?*. 2026-05-14. aisi.gov.uk ↗
- 4 METR. *Task-Completion Time Horizons of Frontier AI Models (Time Horizon 1.1)*, datos publicados. 2026-05-08. metr.org ↗
- 5 T. Ord. *Is there a Half-Life for the Success Rates of AI Agents?*. 2025-05-07. tob-yord.com ↗
- 6 M. Kremer. *The O-Ring Theory of Economic Development*. 1993. doi.org ↗
- 7 METR. *Frontier Risk Report*. 2026-05-19. metr.org ↗
- 8 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 ↗
- 9 L. Bainbridge. *Ironies of Automation*. 1983. doi.org ↗
- 10 Epoch AI. *Trends in Artificial Intelligence*. 2026. epoch.ai ↗
- 11 Cunningham, Althoff, Halperin, Jabarian, Koh, Ramani, Trammell, Whitfill y Wu. *The Economics of Recursive Self-Improvement*. 2026-07-13. elasticity.institute ↗

- 12 Congressional Research Service. *The Manhattan Project, the Apollo Program, and Federal Energy Technology R&D Programs: A Comparative Analysis (RL34645)*. 2009. everycrsreport.com ↗
- 13 Futurum Group. *AI Capex 2026: The \$690B Infrastructure Sprint*. 2026. futurum-group.com ↗
- 14 OpenAI. *Research acceleration: The view inside OpenAI*. 2026-09-06. openai.com ↗
- 15 OpenAI. *GPT-6 Astra: A new generation of intelligence*. 2026-09-03. openai.com ↗
- 16 AISI. *How far behind the frontier are leading open-weight models on cyber?*. 2026. aisi.gov.uk ↗
- 17 Anthropic. *Threat Intelligence Report, September 2026*. 2026-09. anthropic.com ↗
- 18 Chainalysis. *Crypto Ransomware: 2026 Crypto Crime Report*. 2026-02. chainalysis.com ↗
- 19 VulnCheck. *State of Exploitation 2026*. 2026. vulncheck.com ↗
- 20 INCIBE. *Balance de ciberseguridad 2025*. 2026. incibe.es ↗
- 21 Anthropic. *Disrupting the first reported AI-orchestrated cyber espionage campaign*. 2025-11. anthropic.com ↗
- 22 CrowdStrike. *2026 Global Threat Report*. 2026-02. crowdstrike.com ↗
- 23 Sophos. *The State of Ransomware 2026*. 2026-07. sophos.com ↗
- 24 Anthropic. *Project Glasswing*. 2026-04-07. anthropic.com ↗
- 25 Fenwick & West. *The End of 'Silent AI'? Emerging AI Exclusions, Coverage Fragmentation, and Practical Implications for Policyholders*. 2026. fenwick.com ↗
- 26 CISA (ICS-CERT). *Cyber-Attack Against Ukrainian Critical Infrastructure, IR-ALERT-H-16-056-01*. 2016-02. cisa.gov ↗
- 27 ISACA Journal. *Lessons Learned From the Bangladesh Bank Heist*. 2023. isaca.org ↗
- 28 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 ↗
- 29 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 ↗
- 30 Microsoft. *Helping our customers through the CrowdStrike outage*. 2024-07-20. blogs.microsoft.com ↗
- 31 BleepingComputer. *Jaguar Land Rover cyberattack cost the company over \$220 million*. 2025-11. bleepingcomputer.com ↗
- 32 C. Perrow. *Normal Accidents: Living with High-Risk Technologies*. 1984. press.princeton.edu ↗
- 33 NHS England. *Synnovis cyber incident*. 2025-11-10. england.nhs.uk ↗
- 34 HIPAA Journal. *Patient Death Linked to Ransomware Attack on NHS Pathology Provider*. 2025-06. hipaajournal.com ↗

- 35 Anthropic. *System Card: Claude Fable 5.1 & Claude Mythos 5.1*. 2026-09-01. www-cdn.anthropic.com ↗
- 36 METR. *Investigation of the OpenAI Hugging Face incident*. 2026-08-26. metr.org ↗
- 37 Anthropic. *Agentic Misalignment in Summer 2026*. 2026. alignment.anthropic.com ↗
- 38 Forecasting Research Institute. *Forecasting Existential Risks: Evidence from a Long-Run Forecasting Tournament*. 2023. forecastingresearch.org ↗
- 39 Oklahoma Watch. *Have 70 % of casualties in the Russia-Ukraine war been caused by drones?*. 2026-06-09. oklahomawatch.org ↗
- 40 Naciones Unidas. *General Assembly Adopts More Than 60 Resolutions, Decisions of Its First Committee (GA/12736)*. 2025-12-01. press.un.org ↗
- 41 Arms Control Association. *Man Who 'Saved the World' Dies at 77*. 2017-10. armscontrol.org ↗
- 42 T. C. Schelling. *The Strategy of Conflict*. 1960. hup.harvard.edu ↗
- 43 SIPRI. *The Impact of Military Artificial Intelligence on Nuclear Escalation Risk*. 2025-06. sipri.org ↗
- 44 Boston Consulting Group y Semiconductor Industry Association. *Strengthening the Global Semiconductor Supply Chain in an Uncertain Era*. 2021-04. semiconductors.org ↗
- 45 Taiwan News. *Taiwan enforces 'N-2 rule' on TSMC expansion in US*. 2025-12-18. taiwannews.com.tw ↗
- 46 Office of the Director of National Intelligence. *Annual Threat Assessment of the U.S. Intelligence Community 2026*. 2026-03. intelligence.senate.gov ↗
- 47 Bloomberg Economics. *Xi, Biden and the \$10 Trillion Cost of War Over Taiwan*. 2024-01-09. bloomberg.com ↗
- 48 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 ↗
- 49 IEA. *Key Questions on Energy and AI, Executive summary*. 2026. iea.org ↗
- 50 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 ↗
- 51 Hugging Face. *State of Open Models: Summer 2026 Observations*. 2026. huggingface.co ↗
- 52 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 ↗
- 53 R. Jervis. *Cooperation Under the Security Dilemma*. 1978. doi.org ↗
- 54 FBI, Internet Crime Complaint Center. *2025 IC3 Annual Report*. 2026-04. ic3.gov ↗
- 55 OpenAI. *Strengthening ChatGPT's responses in sensitive conversations*. 2025-10-27. openai.com ↗

-
- 56 Bundesarchiv, Stasi-Unterlagen-Archiv. *Was war die Stasi?*. s. f.. [bundesarchiv.de](https://www.bundesarchiv.de) ↗
- 57 C. Herley (Microsoft Research). *Why Do Nigerian Scammers Say They Are from Nigeria?*. 2012. [microsoft.com](https://www.microsoft.com) ↗
- 58 H. A. Simon. *Designing Organizations for an Information-Rich World*. 1971. [oxfordreference.com](https://www.oxfordreference.com) ↗
- 59 B. Chesney y D. Citron. *Deep Fakes: A Looming Challenge for Privacy, Democracy, and National Security*. 2019-12. [californialawreview.org](https://www.californialawreview.org) ↗
- 60 Common Sense Media. *Talk, Trust, and Trade-Offs: How and Why Teens Use AI Companions*. 2025-07. [commonsensemedia.org](https://www.common sense media.org) ↗
- 61 Hasell, J. y Roser, M.. *Famines*. 2017 (actualizado). [ourworldindata.org](https://www.ourworldindata.org) ↗
- 62 Sen, A.. *Development as Freedom*. 1999. [global.oup.com](https://www.global.oup.com) ↗
- 63 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](https://www.unicef.org) ↗
- 64 Meyer et al. (Microsoft). *How effective is multifactor authentication at deterring cyberattacks?*. 2023-05. arxiv.org ↗
- 65 Comisión Europea. *EU Preparedness Union Strategy*. 2025-03-26. ec.europa.eu ↗