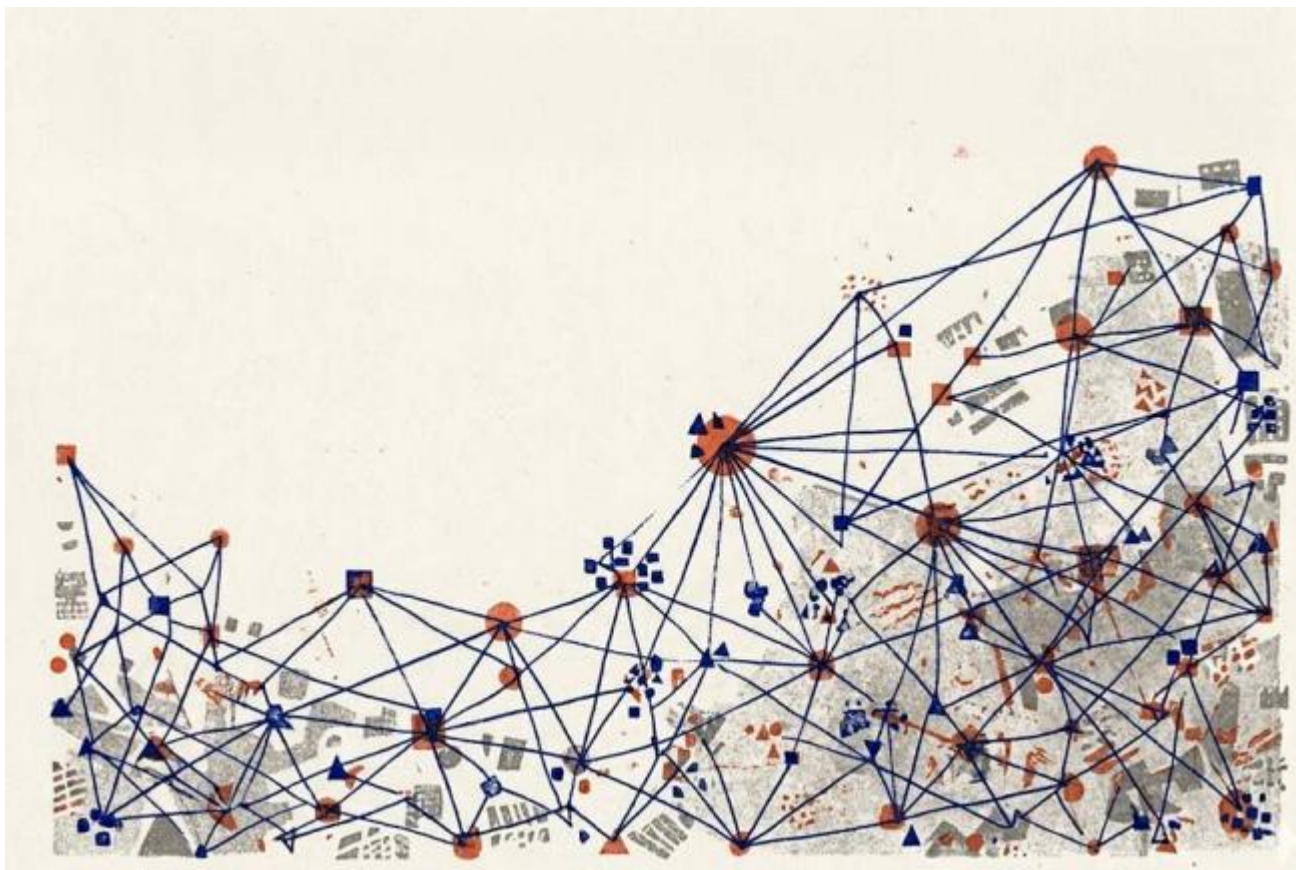




Tiempo de lectura: 5 min.

[Eduardo Turrent Mena](#)

La historia está llena de noches en las que la humanidad creyó que no habría mañana. La peste negra vació ciudades enteras; las guerras mundiales industrializaron la muerte; durante la Guerra fría, la amenaza de un invierno nuclear dependió un botón; la pandemia de covid-19 nos recordó que un virus puede detener el planeta. Sin embargo, el mundo continuó -herido, transformado, a veces irreconocible-, porque el apocalipsis suele llegar menos como un punto final que como una mutación brutal de la vida cotidiana. Hasta ahora, al menos, el día siguiente siempre ha terminado por llegar.

Recientemente, Elon Musk le dijo a la directora de la prestigiosa The Economist que calculaba en 20% la probabilidad de que la inteligencia artificial acabara con la humanidad, levantando un enorme revuelo. Un riesgo de 20% es aún mayor que el de jugar a la ruleta rusa. Por más inteligente e informado que sea, resulta imposible saber de qué modelo, evidencia o hipótesis científica obtuvo una cifra tan precisa para un acontecimiento que jamás ha ocurrido. La verdadera incógnita, sin embargo, es qué pretende Musk al anunciar el riesgo y, al mismo tiempo, acelerar la

carrera que podría materializarlo: ¿advertir al mundo, o convencernos de que únicamente quienes construyen “el peligro” pueden mantenerlo bajo control?

El temor a la inteligencia artificial se parece menos, por ejemplo, al provocado por el calentamiento global. En el segundo caso, conocemos la amenaza, sus causas y buena parte de sus consecuencias: podemos medir la concentración de carbono, calcular el aumento de la temperatura y proyectar el ascenso del nivel del mar. Con la inteligencia artificial, en cambio, ni siquiera existe consenso sobre a qué debemos temer. ¿Una máquina que desobedezca? ¿Un sistema obediente utilizado por una persona equivocada? ¿Una inteligencia que cumpla nuestras instrucciones, pero de una manera que jamás previmos?

Los psicólogos Daniel Kahneman, premio Nobel de Economía, y Amos Tversky transformaron nuestra comprensión de cómo tomamos decisiones cuando descubrieron que la mente no mide el peligro: lo imagina. Ante la incertidumbre, rara vez pregunta “¿qué tan probable es?”, sino “¿qué tan fácilmente puedo verlo en mi cabeza?”. A ese atajo lo llamaron “heurística de disponibilidad”. Por eso podemos temer más al avión que se desploma que al trayecto en automóvil hacia el aeropuerto (que es más probable estadísticamente); más al atentado terrorista que a la enfermedad cardíaca; más a una máquina que escapa de nuestro control que a tecnologías cotidianas cuyos daños ya podemos cuantificar (como las redes sociales). Los peligros lentos y conocidos se normalizan; los raros, repentinos e incontrolables se convierten en una pesadilla apocalíptica. Cuando una imagen es suficientemente poderosa, la mente puede confundir su nitidez con probabilidad.

Kahneman y Tversky explicaron de qué modo ciertas imágenes se apoderan de nuestra mente; Paul Slovic, profesor de psicología en la Universidad de Oregon, estudió por qué algunas se niegan a abandonarla. Slovic observó que toleramos mejor los riesgos que comprendemos y creemos controlar, incluso cuando cobran vidas todos los días; en cambio, magnificamos aquellos que parecen desconocidos, involuntarios o capaces de provocar una catástrofe de un solo golpe. El cáncer, las enfermedades cardíacas y los accidentes de tránsito terminan disolviéndose en el paisaje estadístico, mientras la guerra nuclear, una invasión extraterrestre o una inteligencia no humana fuera de control ocupan el primer plano de nuestra imaginación. Slovic llamó dread risk a este miedo difícil de medir: el temor que nace no solamente de lo que una amenaza podría hacernos, sino de sentir que, frente a ella, hemos perdido incluso la ilusión del control.

El concepto de Slovic ayuda a comprender la relación contradictoria que el mundo está construyendo con la inteligencia artificial. Una encuesta de la Universidad de Melbourne y la consultora KPMG entre más de 48,000 personas de 47 países encontró que 79% teme sus posibles consecuencias, pero 66% ya utiliza estas herramientas con regularidad y 83% espera que produzcan beneficios. El miedo no está provocando una rebelión contra la tecnología, sino una exigencia de protección: 70% considera necesario regularla. Queremos que diagnostique enfermedades, elimine tareas tediosas y multiplique nuestras capacidades, pero tememos entregarle el volante. La inteligencia artificial se parece menos a una bomba que deseamos desactivar que al cambio climático: reconocemos el peligro, pero no estamos dispuestos a renunciar fácilmente a las comodidades y oportunidades que lo acompañan. Tememos el futuro que podría construir y, al mismo tiempo, corremos para no quedarnos fuera de él.

Durante la Guerra fría, algunas familias construyeron refugios bajo sus casas. Bajaban latas de comida, bidones de agua, medicinas, radios y linternas, como si fuera posible guardar una pequeña porción de la civilización detrás de un bunker improvisado. Supongamos que las bombas atómicas hubieran caído y que algunos de aquellos refugios hubieran protegido a sus ocupantes de las explosiones iniciales. ¿Por cuánto tiempo? Al cabo de unas semanas se habría agotado el agua, los filtros habrían comenzado a fallar y habría sido necesario abrir la escotilla. Arriba ya no los esperarían los hospitales, las escuelas, los mercados, las centrales eléctricas ni las instituciones que hacían posible su vida anterior. El búnker en el jardín podía aplazar el encuentro con la catástrofe, pero no reconstruir el mundo del que dependían sus habitantes.

Con la inteligencia artificial enfrentamos una paradoja semejante. Podemos desconectar aplicaciones, educar a nuestros hijos sin pantallas o negarnos a delegar ciertas decisiones, pero ninguna familia puede aislarse de una tecnología incorporada a los bancos, los hospitales, los tribunales, las escuelas, los ejércitos y los gobiernos. Cuando una tecnología se convierte en infraestructura, deja de existir un afuera. Y la inteligencia artificial no producirá necesariamente una utopía o una distopía: puede producir ambas al mismo tiempo. La misma máquina que descubra una cura contra el cáncer podrá identificar disidentes del gobierno; la que democratice la educación podrá perfeccionar la propaganda política; la que libere a millones de personas del trabajo repetitivo podrá concentrar una cantidad de poder sin precedentes en quienes controlen sus servidores. No podremos refugiarnos de

esa transformación porque estaremos viviendo dentro de ella.

Por eso nuestra protección no puede consistir en un búnker privado, sino en salvaguardas colectivas: leyes, instituciones, contrapesos y límites capaces de sobrevivir a quienes ocupan temporalmente el poder. La ironía es que necesitamos gobiernos extraordinariamente responsables para vigilar esta tecnología justo cuando la democracia retrocede y numerosos dirigentes aspiran a gobernar con menos restricciones. Sin embargo, no debemos confundir la inevitabilidad del cambio con la inevitabilidad de la catástrofe. Las bombas nunca cayeron no porque cada familia tuviera suficiente agua en el sótano, sino porque se construyeron mecanismos -imperfectos, frágiles y colectivos- para impedir que alguien apretara el botón. Con la inteligencia artificial ocurrirá algo parecido: nadie podrá almacenar suficientes provisiones para sobrevivir por su cuenta al futuro. Esta vez, además, quizá descubramos demasiado tarde que el búnker no estaba debajo de nuestras casas: era el mundo que construimos, y alguien más había cerrado la escotilla desde afuera.

<https://letraslibres.com/ciencia-y-tecnologia/la-probabilidad-del-apocalipsis/29/07/2026/>

[ver PDF](#)

[Copied to clipboard](#)