

## **Inteligencia artificial: del miedo al apocalipsis a una cultura de calidad y responsabilidad**

Un reciente artículo de Eduardo Turrent Mena, “**La probabilidad del apocalipsis**”, publicado en *Letras Libres* el 29 de julio de 2026, plantea una reflexión particularmente oportuna sobre el temor que acompaña al desarrollo acelerado de la inteligencia artificial.

Artículo original:

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

Turrent parte de una afirmación de Elon Musk, quien estimó en un 20 % la probabilidad de que la inteligencia artificial termine con la humanidad. El autor señala acertadamente un problema fundamental: ¿cómo se obtiene una cifra aparentemente tan precisa para un acontecimiento que nunca ha ocurrido y para el cual carecemos de una base estadística que permita calcular semejante probabilidad?

La pregunta trasciende la anécdota. Nos obliga a distinguir entre **riesgo, incertidumbre y percepción del riesgo**.

### **Cuando imaginar un peligro no significa poder medirlo**

Turrent recurre a los trabajos de Daniel Kahneman y Amos Tversky sobre la denominada *heurística de disponibilidad*. Frente a acontecimientos inciertos, las personas tendemos a considerar más probable aquello que podemos imaginar fácilmente. Un accidente aéreo espectacular puede producir más temor que miles de accidentes automovilísticos; una inteligencia artificial fuera de control resulta mucho más imaginable que los pequeños errores algorítmicos que pueden producirse diariamente sin ocupar titulares.

También recuerda los trabajos de Paul Slovic sobre el *dread risk*: nuestra tendencia a magnificar peligros desconocidos, involuntarios, incontrolables y capaces de producir consecuencias catastróficas.

Esta reflexión es importante porque buena parte de la discusión contemporánea sobre inteligencia artificial oscila entre dos extremos igualmente poco satisfactorios: **la fascinación tecnológica y el catastrofismo**.

Quizás exista una tercera manera de abordar el problema, particularmente familiar para quienes trabajamos en ciencias veterinarias, laboratorios, inocuidad alimentaria y sistemas de calidad: **la gestión del riesgo**.

### **Una lección que los laboratorios conocen desde hace décadas**

Pensemos en un laboratorio moderno.

Un espectrómetro de masas puede distinguir concentraciones que ningún ser humano podría percibir, procesar señales que ningún analista podría calcular manualmente en tiempos razonables y detectar sustancias en cantidades extraordinariamente pequeñas.

Nadie concluye por ello que deberíamos prescindir del instrumento porque ha superado nuestras capacidades.

Pero tampoco aceptamos automáticamente cualquier número que aparezca en su pantalla.

Hemos construido alrededor de esos instrumentos todo un sistema de confianza: métodos validados, calibraciones, materiales de referencia, controles de calidad, estimación de incertidumbre, trazabilidad metrológica, ensayos de aptitud, procedimientos documentados y personal técnicamente competente.

Y, finalmente, alguien asume la responsabilidad por el resultado que el laboratorio informa.

Esta experiencia ofrece una analogía particularmente útil para pensar la inteligencia artificial.

### **El ser humano no necesita ser más inteligente que la IA para supervisarla**

Existe una idea frecuente según la cual el denominado “control humano” significaría que una persona debería ser capaz de comprender, reproducir o superar todo aquello que realiza la inteligencia artificial.

Probablemente esto sea imposible y, además, innecesario.

Ya utilizamos innumerables tecnologías que realizan tareas que superan ampliamente nuestras capacidades individuales. Precisamente por eso las utilizamos.

La responsabilidad humana no exige superioridad intelectual sobre la herramienta. Exige algo diferente: **definir para qué se utiliza, conocer razonablemente sus limitaciones, evaluar críticamente sus resultados, establecer controles adecuados y asumir la responsabilidad por las decisiones adoptadas con su ayuda.**

Esta distinción será cada vez más importante.

La inteligencia artificial puede llegar a revisar en segundos miles de resultados históricos de un laboratorio, identificar tendencias que habían pasado inadvertidas, detectar anomalías, relacionar información procedente de diferentes fuentes o sugerir hipótesis que ningún profesional había considerado.

Ante semejante capacidad, la respuesta responsable no debería ser rechazarla porque “el experto debe saber más que la máquina”.

Debería ser preguntarnos **cómo validamos su utilización.**

### **Del “control humano” a la responsabilidad humana**

Aquí quizá sea necesario introducir un cambio conceptual.

Se suele decir que debemos evitar “entregarle el volante” a la inteligencia artificial. La metáfora es intuitiva, pero puede resultar engañosa.

En determinadas actividades tendremos que permitir que sistemas de IA ejecuten procesos que ningún ser humano podría realizar con la misma velocidad, extensión o complejidad.

La cuestión fundamental no será necesariamente quién tiene físicamente el volante, sino:

**¿Quién determina el destino? ¿Quién establece las reglas? ¿Cómo comprobamos que el sistema funciona dentro de límites aceptables? ¿Quién puede detenerlo? ¿Y quién responde por las consecuencias de utilizarlo?**

Eso es gobernanza.

### **El riesgo de utilizar IA y el riesgo de no utilizarla**

Existe además una dimensión que merece mayor atención.

Toda nueva tecnología introduce riesgos. Pero **renunciar a utilizar una tecnología también puede producirlos.**

Si una herramienta de inteligencia artificial demostrara que puede detectar antes un brote epidemiológico, identificar interacciones desconocidas, descubrir anomalías en resultados analíticos o reducir determinados errores diagnósticos, decidir no utilizarla dejaría de ser una decisión neutral.

La evaluación debería considerar ambos lados:

**riesgo de utilizar IA ↔ riesgo de no utilizar IA.**

Este equilibrio resulta particularmente importante para la regulación.

Necesitamos normas que protejan a las personas y a la sociedad frente a usos irresponsables de la inteligencia artificial. Pero una regulación basada exclusivamente en el temor podría dificultar aplicaciones que aumenten la seguridad, la productividad, la capacidad científica y hasta nuestra posibilidad de detectar riesgos que hoy desconocemos.

### **De la regulación a una cultura de calidad para la IA**

Turrent concluye que, frente a una tecnología que terminará integrada en bancos, hospitales, tribunales, escuelas, ejércitos y gobiernos, ningún individuo podrá construir su propio “búnker”. Necesitaremos salvaguardas colectivas: leyes, instituciones, contrapesos y límites.

Es difícil discrepar con esa conclusión.

Pero quizá podamos agregar una capa adicional.

**Las leyes son necesarias, pero no suficientes.**

También necesitamos desarrollar una verdadera **cultura de calidad en la utilización de la inteligencia artificial.**

Quienes trabajamos con sistemas de calidad sabemos que la confianza no procede simplemente de declarar que algo es seguro. Se construye mediante evidencia, validación, documentación, trazabilidad, controles, auditorías, comparación de resultados y definición de responsabilidades.

¿Por qué no comenzar a trasladar esos mismos principios a la inteligencia artificial?

No significa aplicar literalmente las normas de un laboratorio a un algoritmo. Significa aprovechar una filosofía que disciplinas como la metrología, la industria farmacéutica, la aviación, la medicina veterinaria y la inocuidad alimentaria llevan décadas perfeccionando:

**no confiar ciegamente, pero tampoco renunciar a una herramienta poderosa por miedo a equivocarnos. Validar, controlar, aprender y mejorar.**

### **Después del apocalipsis**

El debate sobre si existe un 10 %, 20 % o 50 % de probabilidad de que una inteligencia artificial termine algún día con la humanidad puede ser intelectualmente fascinante. Pero esos porcentajes no deberían distraernos de decisiones mucho más inmediatas.

La IA ya está entrando en nuestros laboratorios, universidades, empresas, administraciones y profesiones.

Por eso quizá la pregunta relevante no sea:

**¿Qué probabilidad existe de que la inteligencia artificial provoque el apocalipsis?**

Sino:

**¿Cómo aprovechamos una capacidad tecnológica que puede llegar a superar ampliamente nuestras capacidades individuales, manteniendo sistemas verificables de validación, trazabilidad, control y responsabilidad humana?**

La humanidad lleva mucho tiempo aprendiendo a convivir con tecnologías capaces simultáneamente de beneficiarla y perjudicarla.

Con la inteligencia artificial tendremos que volver a hacerlo.

Esta vez, sin embargo, contamos con una ventaja: no partimos de cero. Décadas de experiencia en ciencia, calidad y gestión del riesgo nos han enseñado un principio que puede resultar mucho más útil que intentar adivinar la probabilidad del apocalipsis:

**cuando una tecnología supera nuestras capacidades, la respuesta no es necesariamente utilizarla menos. Es aprender a utilizarla mejor y responsabilizarnos de cómo lo hacemos.**

## Referencia

Eduardo Turrent Mena, “La probabilidad del apocalipsis”, *Letras Libres*, 29 de julio de 2026.

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