

La Ley Europea de Inteligencia Artificial: entre la innovación, la responsabilidad y la incertidumbre

Una lectura crítica de la legislación europea y de las reflexiones de Thierry Breton

La entrada progresiva en aplicación de la Ley de Inteligencia Artificial de la Unión Europea —el conocido *AI Act*— constituye un acontecimiento que trasciende ampliamente el ámbito tecnológico. Estamos ante uno de los primeros intentos sistemáticos de una gran jurisdicción por establecer reglas generales para una tecnología destinada a intervenir crecientemente en la actividad profesional, científica, empresarial y administrativa.

El debate se ha polarizado con facilidad. Para algunos, Europa estaría construyendo una barrera regulatoria que podría perjudicar su capacidad de competir con Estados Unidos y China. Para otros, la Unión Europea está haciendo precisamente aquello que históricamente ha caracterizado su modelo: permitir el desarrollo económico y tecnológico dentro de un marco de protección de derechos, seguridad y responsabilidad.

En este contexto resulta particularmente interesante el reciente artículo de Thierry Breton, excomisario europeo de Mercado Interior y uno de los protagonistas políticos de la elaboración de la legislación europea sobre inteligencia artificial, publicado por *Le Grand Continent* bajo el título “¿Por qué no hemos entendido la Ley IA?”.

Su reflexión proporciona una buena oportunidad no solamente para analizar qué pretende hacer Europa, sino también para discutir algunas de las limitaciones de su estrategia.

Una idea fundamental: la IA no elimina la responsabilidad humana

Probablemente el principio más valioso que puede extraerse tanto de la legislación europea como del razonamiento de Breton es relativamente sencillo: **la inteligencia artificial puede aumentar extraordinariamente las capacidades humanas sin que ello implique transferir a la máquina la responsabilidad por las decisiones adoptadas.**

Esta cuestión adquiere especial importancia en las profesiones científicas y técnicas.

Un médico, veterinario, químico, abogado, ingeniero, investigador o profesional de laboratorio puede utilizar sistemas de inteligencia artificial para buscar información, comparar antecedentes, analizar datos, identificar anomalías, preparar informes o incluso formular recomendaciones.

Pero la existencia de esas herramientas no debería producir una desaparición de la responsabilidad profesional.

Podríamos expresarlo mediante un principio general:

IA para aumentar la competencia profesional; nunca IA para diluir la responsabilidad profesional.

No se trata de impedir que los profesionales utilicen inteligencia artificial. Probablemente ocurrirá exactamente lo contrario: dentro de pocos años resultará difícil imaginar determinadas profesiones científicas sin herramientas avanzadas de IA.

La cuestión fundamental será determinar **quién conserva la capacidad de supervisión y quién responde por la decisión finalmente adoptada.**

Regular el riesgo y no simplemente la tecnología

Una de las características conceptualmente más interesantes del modelo europeo es que intenta clasificar las aplicaciones de inteligencia artificial en función del riesgo asociado a su utilización.

No todas las aplicaciones merecen el mismo tratamiento. Un sistema utilizado para ordenar documentos administrativos no puede evaluarse del mismo modo que otro empleado para seleccionar trabajadores, conceder créditos, determinar determinadas prestaciones públicas o intervenir en decisiones relacionadas con la salud.

La tecnología subyacente puede incluso ser semejante. Lo que cambia radicalmente son **las consecuencias de las decisiones que se adoptan con ella.**

De allí surge la arquitectura de regulación basada en riesgo del AI Act: determinadas prácticas se consideran incompatibles con principios fundamentales y son prohibidas; determinados sistemas considerados de alto riesgo quedan sujetos a obligaciones especialmente rigurosas; otras aplicaciones requieren fundamentalmente transparencia, mientras que una enorme cantidad de usos cotidianos permanece sometida a requisitos mucho menores.

Esta filosofía contiene una enseñanza particularmente importante para las pequeñas y medianas empresas: **no debería regularse una herramienta por el simple hecho de utilizar inteligencia artificial, sino fundamentalmente por aquello que se le permite hacer y por las consecuencias potenciales de sus decisiones.**

Error, equivocación e incertidumbre: una distinción necesaria

Aquí resulta conveniente introducir una precisión científica que frecuentemente desaparece del debate público sobre inteligencia artificial.

Se habla constantemente del “error de la IA”, utilizando la palabra *error* en su significado cotidiano de equivocación. Sin embargo, para las ciencias de la medición conviene distinguir conceptos diferentes.

En metrología, el **error de medida** expresa la diferencia entre un valor medido y un valor de referencia. No significa necesariamente que alguien haya cometido una equivocación.

La **incertidumbre de medida**, por su parte, caracteriza la dispersión de los valores que razonablemente pueden atribuirse al mensurando a partir de la información disponible. Expresa, en términos sencillos, los límites dentro de los cuales conocemos razonablemente un resultado.

Una **equivocación o fallo**, en cambio, corresponde mejor a aquello que en el lenguaje cotidiano llamamos cometer un error: una actuación incorrecta o evitable de una persona, procedimiento o sistema. La distinción no es meramente semántica.

Un sistema de inteligencia artificial puede funcionar exactamente de acuerdo con su diseño y, sin embargo, producir una conclusión caracterizada por una elevada incertidumbre porque los datos disponibles son insuficientes, variables, contradictorios o ambiguos.

Por ello, uno de los grandes desafíos de la inteligencia artificial científica será conseguir que las máquinas no solamente produzcan respuestas, sino que sean capaces de **reconocer, estimar y comunicar adecuadamente los límites de aquello que saben**.

La tecnología subyacente puede ser semejante. La incertidumbre asociada a sus resultados y, sobre todo, las consecuencias de una decisión incorrecta pueden ser completamente diferentes.

Esta cultura de la incertidumbre, profundamente incorporada a la metrología y a las ciencias experimentales, debería ocupar un lugar mucho más importante en la futura gobernanza de la inteligencia artificial.

¿Ha frenado Europa la innovación?

En este punto el análisis de Thierry Breton resulta más discutible.

El excomisario rechaza la idea de que el AI Act sea responsable del retraso europeo frente a Estados Unidos en el desarrollo de grandes empresas de inteligencia artificial. Su argumento tiene una base razonable: buena parte de la extraordinaria expansión estadounidense de los últimos años ocurrió cuando muchas de las obligaciones europeas todavía no eran aplicables.

Breton señala otros problemas estructurales europeos: fragmentación de los mercados de capitales, menor disponibilidad de capital de riesgo, dificultades para financiar empresas tecnológicas durante sus fases de expansión y una cultura empresarial comparativamente más adversa al riesgo.

Es difícil discutir la importancia de estos factores. Sin embargo, afirmar que una legislación que todavía no ha entrado plenamente en aplicación no puede afectar a las decisiones empresariales sería ir demasiado lejos.

Las inversiones se realizan mirando hacia el futuro. Un empresario o inversor que prevé que dentro de algunos años necesitará procedimientos adicionales de documentación, evaluación, auditoría, trazabilidad o asesoramiento jurídico puede incorporar esos costes esperados a su decisión de inversión actual.

Existe además otro problema específicamente europeo: **la acumulación regulatoria**. Cada regulación puede tener una justificación perfectamente razonable cuando se considera individualmente. Pero la superposición de normas sobre inteligencia artificial, datos, privacidad, ciberseguridad, servicios digitales y legislación sectorial puede producir conjuntamente costes de cumplimiento considerables.

Para una gran corporación internacional esos costes pueden representar una partida adicional de su estructura jurídica y regulatoria. Para una pequeña empresa innovadora pueden convertirse en una barrera de entrada.

Europa deberá vigilar cuidadosamente esta diferencia.

“Bruselas” no es un legislador extranjero

Existe además una dimensión política que frecuentemente queda escondida detrás del lenguaje utilizado para criticar las regulaciones europeas. Se dice habitualmente que “Bruselas impone” determinadas normas.

Pero el AI Act fue aprobado por el Parlamento Europeo por una amplísima mayoría y posteriormente respaldado por los Estados miembros mediante los mecanismos institucionales de la Unión.

Esto obliga a introducir una consideración democrática. Los ciudadanos europeos pueden considerar una regulación acertada, insuficiente o excesiva. Pueden exigir su modificación y votar a representantes que defiendan políticas diferentes.

Pero resulta intelectualmente problemático presentar permanentemente a “Bruselas” como si fuese una potencia extranjera que legislara sobre los europeos.

Las instituciones comunitarias forman parte del sistema político que los propios ciudadanos europeos contribuyen a constituir mediante sus elecciones nacionales y europeas. La responsabilidad regulatoria es, por tanto, también una **responsabilidad democrática**.

Estados Unidos, China y Europa: tres modelos diferentes

El artículo de Breton adquiere una dimensión particularmente interesante cuando sitúa el AI Act dentro de la competencia tecnológica mundial.

Simplificando inevitablemente una realidad mucho más compleja, pueden identificarse tres grandes aproximaciones.

Estados Unidos ha favorecido históricamente un modelo en el cual la innovación empresarial dispone de un espacio considerable y muchos de los problemas son abordados posteriormente mediante regulación, litigación o intervención administrativa.

Europa se inclina tradicionalmente hacia un modelo precautorio: intenta establecer previamente determinadas garantías para que la innovación se desarrolle dentro de límites jurídicos.

China introduce un tercer elemento: una política industrial en la cual las capacidades tecnológicas privadas y públicas pueden integrarse dentro de objetivos estratégicos nacionales y recibir diferentes formas de apoyo estatal.

Por ello, la competencia mundial en inteligencia artificial no debería interpretarse simplemente como una carrera entre empresas.

Estamos observando también una competencia entre **modelos económicos, regulatorios y geopolíticos**.

El poder regulatorio europeo y sus límites

Europa posee una ventaja extraordinaria: un mercado de aproximadamente 450 millones de habitantes, con elevada capacidad adquisitiva y un marco jurídico integrado. Una empresa global difícilmente puede renunciar a ese mercado.

De allí surge el denominado *Brussels Effect*: determinadas normas europeas terminan influyendo sobre estándares internacionales porque las empresas encuentran más eficiente adaptar globalmente sus productos a requisitos europeos que mantener sistemas completamente diferentes para cada mercado.

El Reglamento General de Protección de Datos constituye probablemente el ejemplo más conocido. Breton considera que algo semejante podría ocurrir con la inteligencia artificial. Es perfectamente posible.

Pero existe una diferencia importante entre **tener capacidad para regular una tecnología y tener capacidad para producirla**.

Europa dispone de mercado, instituciones científicas extraordinarias y considerable poder regulatorio. Sin embargo, continúa teniendo menor presencia que Estados Unidos en algunos componentes fundamentales del ecosistema de inteligencia artificial: grandes plataformas digitales, proveedores globales de infraestructura de nube, capital de riesgo a gran escala y modelos fundacionales líderes.

Por ello, el objetivo europeo no debería consistir en reducir su soberanía regulatoria.

Debería ser más ambicioso: **añadir soberanía tecnológica a la soberanía regulatoria**.

Regular tecnologías fundamentalmente producidas por otros puede proteger a los ciudadanos europeos. Producirlas además de regularlas genera también poder económico y geopolítico.

El próximo desafío: los agentes de inteligencia artificial

Hay una observación de Breton particularmente significativa: reconoce que será necesario desarrollar un marco adecuado para los **agentes autónomos**, fenómeno cuya importancia era mucho menor cuando comenzó a diseñarse el AI Act.

Estamos entrando rápidamente en una nueva etapa. La primera generación de inteligencia artificial generativa fundamentalmente respondía preguntas. La siguiente comenzó a razonar, analizar documentos y utilizar herramientas.

Ahora aparecen sistemas capaces de **ejecutar acciones**. La diferencia es fundamental.

Una IA que proporciona una respuesta incorrecta genera fundamentalmente un problema de información.

Un agente capaz de modificar una base de datos, emitir una comunicación, realizar una compra, activar un procedimiento, gestionar otro programa informático o intervenir en una decisión profesional plantea además un problema de **acción y responsabilidad**.

Esto obligará probablemente a desarrollar una segunda generación de gobernanza de la inteligencia artificial.

La pregunta dejará progresivamente de ser solamente:

¿Qué puede decir una IA?

para convertirse en:

¿Qué estamos dispuestos a permitir que haga una IA por nosotros?

Tres niveles para incorporar IA en las organizaciones

Esta evolución permite proponer una clasificación sencilla que puede resultar especialmente útil para empresas, laboratorios, universidades, administraciones y organizaciones profesionales.

Un primer nivel sería la **IA asistente**: busca información, resume, compara documentos, redacta borradores y organiza conocimiento.

Un segundo nivel sería la **IA supervisada**: analiza datos, detecta anomalías, formula recomendaciones y prepara decisiones que posteriormente revisa una persona competente.

El tercer nivel sería la **IA agente**: recibe autorización para ejecutar determinadas acciones dentro de límites previamente definidos.

La progresión entre estos tres niveles debería ir acompañada de un aumento paralelo de los requisitos de trazabilidad, validación, supervisión y responsabilidad.

No necesitamos necesariamente la misma intensidad regulatoria para una IA que resume veinte artículos científicos que para otra autorizada a intervenir directamente en un procedimiento que pueda afectar derechos, salud, seguridad o patrimonio.

¿Existe espacio para una vía iberoamericana?

El debate europeo contiene finalmente una enseñanza particularmente relevante para Iberoamérica.

Nuestros países no necesitan copiar automáticamente ninguno de los grandes modelos.

Adoptar completamente el enfoque estadounidense puede resultar problemático porque muchos Estados iberoamericanos poseen menor capacidad institucional para corregir posteriormente los efectos adversos de determinadas tecnologías.

Copiar literalmente toda la arquitectura regulatoria europea también podría resultar inconveniente. Una proporción muy elevada del tejido empresarial iberoamericano está constituida por pequeñas y medianas empresas, para las cuales los costes regulatorios relativos pueden ser considerablemente mayores.

Y adoptar tecnologías procedentes de China simplemente porque sean económicas tampoco constituye por sí mismo una estrategia de soberanía tecnológica.

Existe posiblemente espacio para desarrollar una **cuarta vía iberoamericana**.

Podría combinar la protección europea de los derechos fundamentales con una regulación estrictamente proporcional al riesgo, una clara identificación de la responsabilidad humana y una amplia libertad de experimentación para aplicaciones de bajo riesgo.

Su principio podría resumirse así: **regular las consecuencias sin obstaculizar innecesariamente la herramienta**.

De la cultura del error a la cultura de la incertidumbre

Quizá las ciencias experimentales puedan aportar al debate sobre inteligencia artificial algo que la discusión jurídica y tecnológica todavía no ha incorporado suficientemente.

La ciencia moderna no promete certeza absoluta. Aprendió, por el contrario, a medir, declarar y gestionar incertidumbres. Ésa puede convertirse en una enseñanza extraordinariamente valiosa para la inteligencia artificial.

No deberíamos exigir que una IA sea infalible. Ningún profesional humano lo es y ninguna metodología científica proporciona conocimiento absolutamente libre de incertidumbre.

Deberíamos exigir algo más realista y científicamente mucho más sólido: que podamos conocer razonablemente **hasta dónde podemos confiar en un resultado, qué incertidumbre lo acompaña, qué información sustentó la conclusión y quién asume finalmente la responsabilidad por actuar sobre ella**.

Desde esta perspectiva, la futura regulación de la inteligencia artificial no debería limitarse a determinar quién responde cuando un sistema se equivoca.

Debería establecer también **quién evalúa, comunica y asume la incertidumbre cuando una inteligencia artificial participa en una decisión que afecta a otras personas**.

Europa ha dado un primer paso extraordinariamente importante al intentar construir una gobernanza de la inteligencia artificial basada en el riesgo.

El desafío que comienza ahora es mucho más complejo: demostrar que es posible proteger derechos y exigir responsabilidad sin transformar el cumplimiento regulatorio en una ventaja reservada a las grandes corporaciones capaces de pagarlo.

Y para Iberoamérica existe una oportunidad adicional: observar críticamente las experiencias estadounidense, china y europea antes de construir nuestras propias reglas.

No se trata de elegir entre innovación y regulación.

Se trata de encontrar una arquitectura institucional en la cual **innovación, responsabilidad, conocimiento de la incertidumbre y libertad empresarial puedan coexistir.**

Referencia principal: Thierry Breton, “¿Por qué no hemos entendido la Ley IA?”, *Le Grand Continent*, 13 de agosto de 2026. El presente artículo analiza tanto los planteamientos del autor como el enfoque general de la legislación europea sobre inteligencia artificial, incorporando una valoración crítica independiente.