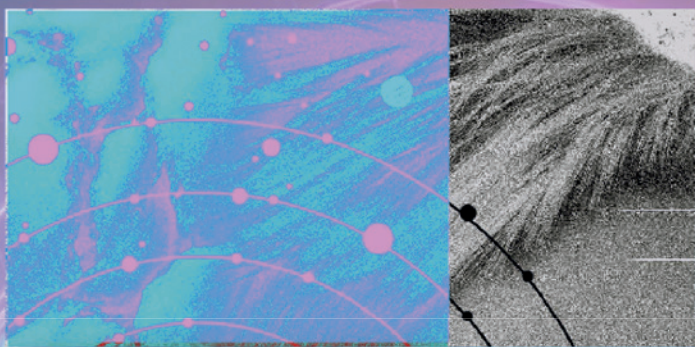


EDICIÓN
2024

PANORAMA



INTELIGENCIA ARTIFICIAL: INFORMACIÓN, SENSIBILIDAD Y ÉTICA PARA UNA CIUDADANÍA CONSCIENTE



INTELIGENCIA ARTIFICIAL, ÉTICA Y CIUDADANÍA CONSCIENTE



unesco

Miembro de
asociaciones y clubes

Asociación
Gran Canaria
para la UNESCO



unesco
Federación Española de
Asociaciones y Clubes
para la UNESCO



CONSEJERÍA
DE PRESIDENCIA
Y MOVILIDAD SOSTENIBLE



ÉTICA Y CIUDADANÍA CONSCIENTE

La integración de la inteligencia artificial (IA) en la educación para la ciudadanía debe fomentar la capacidad crítica y participativa. Esto implica no solo proporcionar conocimientos sino, también, habilidades para analizar, cuestionar y formar opiniones informadas.

Las instituciones educativas, los medios de comunicación y las organizaciones civiles jugamos un papel fundamental en la comprensión de cómo la IA afecta a la vida cotidiana, en la participación en el debate sobre sus regulaciones y en la toma de decisiones responsables respecto a su uso.

Al estar informados y comprometidos, la ciudadanía puede garantizar que la IA se utilice de manera ética y en beneficio del bienestar social, protegiendo los derechos fundamentales y promoviendo una sociedad más equitativa.

La Asociación Gran Canaria para la UNESCO (AGCU), en colaboración con el Cabildo de Gran Canaria, a través de su proyecto 'Inteligencia Artificial: Información, Sensibilidad y Ética para una Ciudadanía Consciente', asume el desafío de contribuir a que la IA sirva para enriquecer y no para debilitar la participación ciudadana.

Para ello, la AGCU prevé la edición anual del boletín digital 'PANORAMA IA', cuyo primer número aborda los aspectos y desafíos éticos que presentan las tecnologías de IA. La 'ética de la IA' es un área clave que se ocupa de las implicaciones morales y sociales del desarrollo y uso de sistemas de IA. Por ello, esta primera entrega trata de crear conciencia sobre esta cuestión trascendental, liderada en el mundo por la UNESCO a través de su 'Recomendación sobre la Ética de la Inteligencia Artificial'.

A medida que la IA avanza, su influencia sobre aspectos de la vida cotidiana, el trabajo y la sociedad aumenta, lo que hace necesario establecer marcos éticos que guíen su uso de manera responsable y justa.

Los principios éticos en la IA buscan asegurar que las tecnologías sean transparentes, responsables y respetuosas con los derechos humanos fundamentales.

“Para abordar los dilemas éticos que surgen con la inteligencia artificial, es vital que tecnólogos, filósofos, legisladores y la sociedad civil trabajen juntos. Solo así se podrán crear marcos normativos y establecer principios que guíen el desarrollo y uso responsable de estas tecnologías. Se trata de garantizar que los riesgos se reduzcan y los beneficios se maximicen para todos” ■

Juan José Benítez de Lugo y Massieu,
Presidente, Asociación Gran Canaria para la UNESCO



¿Qué es la inteligencia artificial (IA)?



Le preguntamos a ChatGPT...

ChatGPT (sistema de chat basado en el modelo de lenguaje por Inteligencia Artificial GPT-3.5, desarrollado por la empresa OpenAI) responde a la pregunta:

¿Qué es la IA?

«La **Inteligencia Artificial (IA)** es una rama de la informática que busca desarrollar sistemas capaces de realizar tareas que requieren de inteligencia humana, como el reconocimiento de patrones, la toma de decisiones, el aprendizaje y la resolución de problemas. Estos sistemas simulan procesos cognitivos humanos, como el pensamiento lógico, el razonamiento y la comprensión del lenguaje, con el fin de automatizar y mejorar diversos aspectos de la vida cotidiana y profesional.

El concepto de IA engloba una variedad de tecnologías, desde el **aprendizaje automático** (machine learning), que permite a las máquinas aprender de los datos sin ser programadas explícitamente para cada tarea, hasta la **visión por computadora**, el **procesamiento del lenguaje natural** (NLP), y la **robótica**. Gracias a estas técnicas, la IA se aplica en múltiples campos: desde asistentes virtuales y motores de recomendación, hasta diagnósticos médicos y vehículos autónomos.

A pesar de los avances significativos en la IA, esta enfrenta **retos complejos**. Algunos de ellos incluyen:

1. **Seguridad y ética**: A medida que la IA se vuelve más poderosa, surge la preocupación sobre el control de sus decisiones, la posible discriminación algorítmica y el uso indebido de datos sensibles. Además, el impacto en el empleo debido a la automatización es un desafío social significativo.
2. **Interpretabilidad y transparencia**: Muchos modelos de IA, especialmente los de aprendizaje profundo, funcionan como “cajas negras”, lo que dificulta entender cómo llegan a sus conclusiones. Esto plantea problemas en áreas como la salud o la justicia, donde la transparencia es crucial.
3. **Datos y privacidad**: La IA requiere grandes volúmenes de datos para funcionar, lo que genera preocupaciones sobre la recopilación, uso y protección de la información personal. Garantizar la privacidad y el uso ético de los datos es uno de los principales obstáculos.



4. ****Generalización****: La mayoría de las IA actuales son “estrechas”, es decir, están diseñadas para realizar tareas específicas. El reto de crear una “inteligencia general” que pueda adaptarse a cualquier situación, tal como lo haría un humano, sigue siendo uno de los grandes objetivos de la investigación en este campo.

En resumen, la IA promete transformar profundamente muchos aspectos de la sociedad y la industria, pero también requiere una cuidadosa gestión de sus riesgos y limitaciones para garantizar su desarrollo responsable».■

¿Conoces los dilemas éticos que plantea la inteligencia artificial?

La inteligencia artificial (IA) es una tecnología que afecta cada vez más a nuestras vidas y tiene el potencial de transformar nuestra sociedad y nuestro planeta, para bien y para mal.

La IA plantea dilemas éticos debido a su capacidad para realizar tareas que antes eran exclusivas de los humanos, lo que genera preocupaciones en áreas como la privacidad, la seguridad, la equidad y la autonomía humana.

Los desarrollos recientes en la IA son el resultado del aumento de la potencia de procesamiento, las mejoras en los algoritmos y el crecimiento exponencial en el volumen y la variedad de datos digitales. Las tecnologías de vanguardia de la IA, incluidos los modelos de IA generativa (IAG)* que proliferan rápidamente, son cada vez más potentes y capaces de realizar numerosas tareas.

El rápido auge de la IA ha generado nuevas **oportunidades**, beneficiosas a nivel global, al facilitar los diagnósticos de salud; el desarrollo de medios de transporte más seguros; el aumento de la eficiencia laboral mediante la automatización de tareas; la creación de productos y servicios personalizados, más baratos y duraderos, y ofrecer otras muchas ventajas para la ciudadanía.

Sin embargo, estos rápidos cambios también plantean **amenazas e impactos** generales para la humanidad y estructurales para la sociedad, que pueden afectar de modo importante a los derechos fundamentales y a los principios democráticos y constitucionales.

Los **dilemas éticos** más controvertidos relacionados con la inteligencia artificial, tienen que ver con la **privacidad y la protección de datos**, la **manipulación del comportamiento**, la **supervisión humana**, la **gobernanza de datos**, la **transparencia**, la **diversidad y la justicia**, el **bienestar social y ambiental**, y la **responsabilidad**.

El diseño, desarrollo y uso de sistemas de IA deben sustentarse en principios y valores éticos para asegurar que sean justos, responsables y respetuosos con los derechos humanos. ■



*IA generativa (IAG): Es un tipo de IA que utiliza algoritmos de aprendizaje automático profundo (Deep Learning) para generar nuevos datos o contenidos a partir de los ya existentes, imitando o mejorando su estilo, calidad y estructura.



Los siguientes ejemplos pueden ayudarnos a entender algunos de los importantes dilemas éticos que presenta la IA y que exigen equilibrar los beneficios tecnológicos con las preocupaciones sobre derechos, equidad y justicia:

1. Sesgo algorítmico y discriminación

Dilema: Los sistemas de IA se entrenan con grandes cantidades de datos, pero si esos datos contienen sesgos (por ejemplo, por razones de género, raza o clase social), la IA puede perpetuar o amplificar esas discriminaciones.

Ejemplo: Un algoritmo de selección de personal que prefiera a candidatos hombres porque ha sido entrenado con datos históricos sesgados en los que los hombres eran contratados más frecuentemente.

Preguntas éticas: ¿Cómo evitar que los sistemas de IA reflejen los prejuicios sociales? ¿Cómo corregir los algoritmos sesgados y garantizar decisiones justas?

2. Privacidad y vigilancia

Dilema: Los sistemas de IA pueden analizar enormes cantidades de datos personales, lo que plantea riesgos para la privacidad y el control de la información personal.

Ejemplo: Cámaras de vigilancia con sistemas de reconocimiento facial que identifican a las personas en tiempo real, lo que puede ayudar a mejorar la seguridad pero también podría vulnerar la privacidad y dar lugar a vigilancia masiva sin consentimiento.

Preguntas éticas: ¿Cuánta vigilancia es aceptable en nombre de la seguridad? ¿Cómo equilibrar la necesidad de seguridad con la protección de la privacidad?

3. Decisiones automatizadas y autonomía

Dilema: Los sistemas de IA están tomando decisiones que antes eran exclusivamente humanas, desde aprobaciones de créditos hasta diagnósticos médicos. Aunque pueden ser más eficientes, estas decisiones automatizadas pueden ser difíciles de cuestionar o explicar.

Ejemplo: Un sistema de IA rechaza una solicitud de crédito basada en datos financieros sin ofrecer una explicación clara al solicitante.

Preguntas éticas: ¿Deberían las personas tener derecho a apelar las decisiones tomadas por IA? ¿Cómo garantizar la transparencia y explicabilidad en las decisiones automatizadas?

4. Desplazamiento de empleos por la automatización

Dilema: La IA está reemplazando tareas humanas en muchos sectores, lo que puede aumentar la eficiencia pero también lleva al desempleo de trabajadores, especialmente en sectores industriales o de servicios.

Ejemplo: Las fábricas que usan robots y sistemas de IA para realizar tareas que antes eran realizadas por operarios humanos, lo que lleva a la pérdida de puestos de trabajo.

Preguntas éticas: ¿Cómo apoyar a los trabajadores desplazados por la IA? ¿Debería haber políticas que limiten la automatización en ciertos sectores o que ofrezcan un mayor apoyo a la reconversión laboral?

5. Uso de IA en conflictos militares

Dilema: La IA está siendo utilizada en el desarrollo de armas autónomas que pueden tomar decisiones sobre atacar sin intervención humana. Esto plantea preocupaciones sobre la ética de delegar decisiones de vida o muerte a una máquina.

Ejemplo: Drones militares equipados con IA que pueden identificar y atacar objetivos sin supervisión humana directa.

Preguntas éticas: ¿Debería permitirse que la IA tome decisiones letales? ¿Deberían prohibirse las armas autónomas en los conflictos bélicos?

6. Manipulación a través de IA en redes sociales

Dilema: Los algoritmos de IA de las plataformas de redes sociales determinan qué contenido se muestra a los usuarios, lo que puede influir en su comportamiento, opiniones y decisiones políticas.

Ejemplo: Un algoritmo que promueve contenido sensacionalista o polarizador porque genera más interacción, lo que lleva a la desinformación o a la radicalización de las opiniones.

Preguntas éticas: ¿Qué responsabilidad tienen las plataformas de redes sociales al influir en la información que consumen las personas? ¿Cómo garantizar que los algoritmos no favorezcan la desinformación o la manipulación?

7. IA en el ámbito médico

Dilema: Los sistemas de IA están siendo usados para diagnósticos y tratamientos médicos, lo que mejora la precisión, pero plantea cuestiones sobre la responsabilidad y el rol del médico.

Ejemplo: Un sistema de IA sugiere un tratamiento para un paciente que un médico humano no habría recomendado, y si la decisión del algoritmo resulta perjudicial, no está claro quién es responsable.

Preguntas éticas: ¿Quién es responsable si un diagnóstico de IA es incorrecto? ¿Deberían los médicos confiar plenamente en las recomendaciones de un sistema de IA?

8. Manipulación en marketing

Dilema: La IA puede personalizar las estrategias de marketing y publicidad para ser extremadamente efectivas en influir en los consumidores, pero esto puede limitar la capacidad de tomar decisiones libres e informadas.

Ejemplo: Un sistema de publicidad personalizada que utiliza datos de comportamiento para dirigir ofertas a individuos vulnerables (por ejemplo, personas con problemas de juego) y manipularlos para gastar dinero.

Preguntas éticas: ¿Hasta qué punto es ético influir en las decisiones de los consumidores mediante IA? ¿Cómo garantizar que los usuarios comprendan el grado en que la IA está influenciando sus elecciones?

9. Relación humano-máquina y dependencia

Dilema: A medida que la IA interactúa de manera más natural con los seres humanos, como en el caso de los asistentes virtuales, las personas pueden desarrollar relaciones o dependencias con las máquinas.

Ejemplo: Un asistente de IA que responde a las emociones del usuario, creando una relación percibida como personal y emocional, lo que puede llevar a dependencias emocionales o a la erosión de la interacción humana.

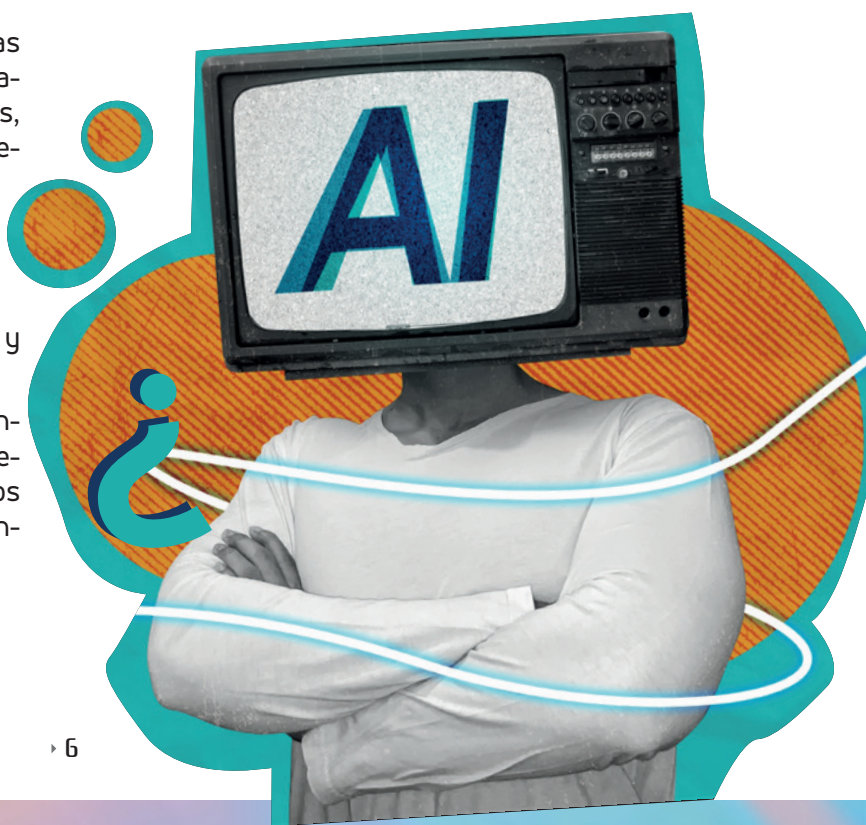
Preguntas éticas: ¿Es ético diseñar IA para fomentar relaciones emocionales con los usuarios? ¿Cómo evitar que las personas dependan excesivamente de las interacciones con IA?

10. IA y Derechos Humanos en regímenes autoritarios

Dilema: En algunos regímenes autoritarios, la IA está siendo utilizada para vigilar y controlar a la población, lo que plantea serias preocupaciones sobre derechos humanos.

Ejemplo: El uso de IA para la vigilancia masiva y el reconocimiento facial en minorías étnicas, como los uigures en China, lo que facilita la represión y la discriminación sistemática.

Preguntas éticas: ¿Cómo deben los desarrolladores y países responder ante el uso de IA para violar los derechos humanos? ¿Deberían existir sanciones internacionales para evitar el uso de IA en contextos de abuso de poder? ■



¿Sabías que la UNESCO lidera el uso regulado y responsable de la inteligencia artificial en el mundo?

La UNESCO desempeña un claro liderazgo en el uso ético y regulado de la IA. Su actividad pionera en este ámbito, le permite guiar el desarrollo y la preparación de aquellos aspectos éticos que la IA precisa con el objetivo de armonizar tecnología y valores humanos en todo el mundo.

La Recomendación de la UNESCO sobre la Ética de la Inteligencia Artificial, adoptada por sus 193 Estados miembros en noviembre de 2021, ha sido el primero y sigue siendo el único marco ético mundial sobre este tema, habiendo logrado avances tangibles en sus dos primeros años de existencia. Más de 50 países participan en la actualidad activamente en la aplicación de este instrumento, y la cooperación multilateral ha aumentado considerablemente.

La Recomendación sobre la Ética de la IA es una contribución fundamental de la UNESCO al objetivo de una gobernanza eficaz y ética de la IA mediante la adopción de una ambiciosa normativa a nivel global.

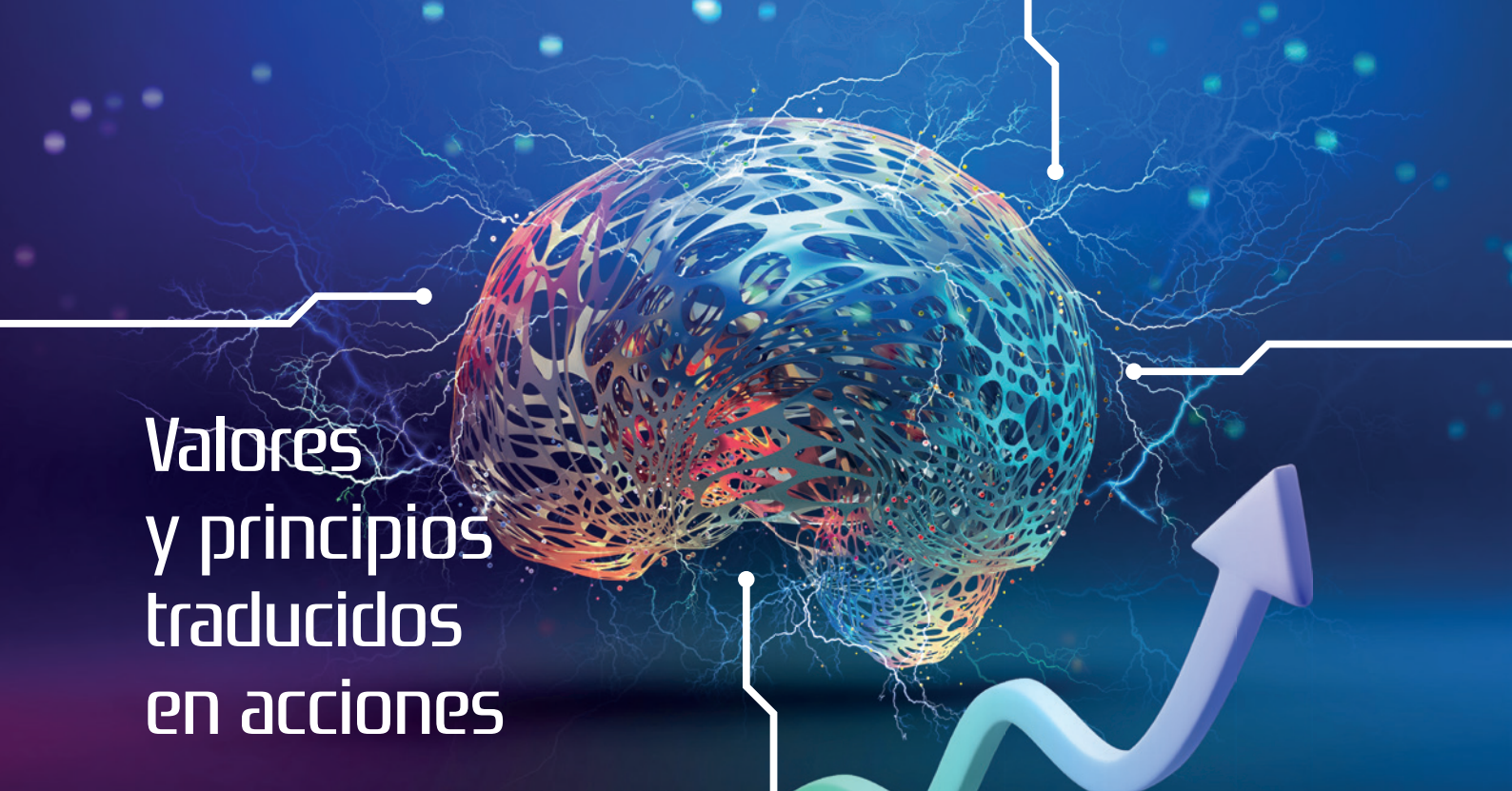
La Recomendación subraya las implicaciones éticas de la inteligencia artificial por su impacto en múltiples ámbitos, como en “la adopción de decisiones, el empleo y el trabajo, la interacción social, la atención de la salud, la educación, los medios de comunicación, el acceso a la información, la brecha digital, la protección del consumidor y de los datos personales, el medio ambiente, la democracia, el estado de derecho, la seguridad y el mantenimiento del orden, el doble uso y los derechos humanos y las libertades fundamentales, incluidas la libertad de expresión, la privacidad y la no discriminación”.

También, llama la atención sobre los desafíos éticos que surgen debido al potencial de los algoritmos de la IA “para reproducir y reforzar los sesgos existentes, lo que puede exacerbar las formas ya existentes de discriminación, los prejuicios y los estereotipos”.

La protección de los derechos humanos y la dignidad es la piedra angular de la Recomendación, basada en principios fundamentales como la transparencia y la equidad, recordando siempre la importancia de la supervisión humana de los sistemas de IA.

La puesta en práctica por parte de los Estados miembros de la ‘Recomendación sobre la Ética de la IA’ puede realizarse a través de herramientas y metodologías innovadoras, como la Metodología de Evaluación del Grado de Preparación y la Evaluación del Impacto Ético.

Mediante la aplicación de estas herramientas, la UNESCO está cambiando el modelo de negocio que impulsa la IA, con vistas a desarrollar una serie de soluciones concretas y prácticas, y garantizar que los resultados de la IA sean justos, inclusivos, sostenibles y no discriminatorios. ■



Valores y principios traducidos en acciones

La Recomendación sobre la Ética de la Inteligencia Artificial de la UNESCO, establece una serie de valores y principios comunes para orientar el desarrollo saludable y responsable de la IA, como el respeto a la diversidad cultural, la igualdad de género, la solidaridad, la justicia, la participación, el bienestar humano y el desarrollo sostenible.

“Los valores y principios [...] deberían ser respetados por todos los actores durante el ciclo de vida de los sistemas de IA, en primer lugar, y, cuando resulte necesario y conveniente, ser promovidos mediante modificaciones de las leyes, los reglamentos y las directrices empresariales existentes y la elaboración de otros nuevos. Todo ello debe ajustarse al derecho internacional, en particular la Carta de las Naciones Unidas y las obligaciones de los Estados Miembros en materia de derechos humanos, y estar en consonancia con los [...] Objetivos de Desarrollo Sostenible (ODS) de las Naciones Unidas”. (Recomendación nº 61)

Lo que permite que la Recomendación sea aplicable, son sus amplios ámbitos de acción política, que permiten a los responsables políticos traducir estos valores y principios fundamentales en acciones con respecto a la gobernanza de datos, el medio ambiente, el género, la educación, la investigación, la salud y el bienestar social, entre otros muchos.

Desde la adopción de la Recomendación, la UNESCO también ha avanzado en la cooperación con el sector privado, lo que ha llevado a la creación de un Consejo Empresarial para la Ética de la IA.■

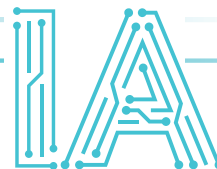
Impulso a la regulación de la ética de la IA por parte de los gobiernos

La Recomendación de la UNESCO propone a los gobiernos la adopción de marcos regulatorios que garanticen que las tecnologías de inteligencia artificial estén al servicio de los intereses de los ciudadanos y beneficien a la humanidad en su conjunto.

Esto se refleja en el nuevo Reglamento de IA de la Unión Europea, cuyo texto final se ha publicado en el Diario Oficial de la Unión Europea el 12 de julio de 2024, cuya relevancia y trascendencia se explica más adelante en este documento.

La ética de la IA no solo depende del cumplimiento de normas y regulaciones sino de actitudes y comportamientos. Para ello, se debe fomentar el pensamiento crítico, el diálogo intercultural, el respeto a los derechos humanos y la diversidad, y la participación ciudadana ■.

PRINCIPIOS ÉTICOS Fundamentales de la



1

TRANSPARENCIA ■

- La IA debe ser comprensible y explicable para los usuarios. Esto significa que los algoritmos y las decisiones que toman deben ser claros, y es fundamental que los ciudadanos entiendan cómo y por qué se toman ciertas decisiones.
- El concepto de “caja negra” de los algoritmos, donde las decisiones son opacas y difíciles de explicar, plantea un desafío ético significativo. Para abordar esto, se está promoviendo la explicabilidad algorítmica.

2

RESPONSABILIDAD ■

- Es crucial definir quién es responsable de las decisiones y acciones tomadas por sistemas de IA, especialmente cuando hay impactos negativos.
- Las empresas, los gobiernos y los desarrolladores deben asumir la responsabilidad legal y ética por los resultados de sus sistemas, tanto si son intencionados como si no lo son.

3

NO DISCRIMINACIÓN Y JUSTICIA ■

- La IA debe evitar reproducir o amplificar sesgos y discriminaciones, como los basados en raza, género, edad o estatus socioeconómico.
- Esto implica que los desarrolladores deben abordar los sesgos algorítmicos mediante un diseño inclusivo y la evaluación continua de los datos que alimentan los sistemas, asegurando que no generen discriminación injusta.

4

PRIVACIDAD Y PROTECCIÓN DE DATOS ■

- La IA puede implicar el tratamiento masivo de datos personales, lo que hace fundamental garantizar la protección de la privacidad de los usuarios.
- El Reglamento General de Protección de Datos (RGPD) de la UE establece que los ciudadanos tienen derecho a saber cómo se utilizan sus datos, a dar su consentimiento explícito y a solicitar la eliminación de sus datos cuando sea necesario.

5

SEGURIDAD ■

- Los sistemas de IA deben ser seguros, fiables y resistentes a ataques cibernéticos o manipulación. Es esencial que la IA funcione de manera controlada y que los posibles errores o problemas sean prevenidos y corregidos rápidamente.
- El control humano sobre las IA críticas (como en la medicina o la seguridad) es otro tema crucial. Los seres humanos deben mantener siempre la capacidad de intervención.

6

AUTONOMÍA ■

- Uno de los debates clave en la ética de la IA es el equilibrio entre la autonomía de las máquinas y el control humano. La IA no debe privar a los seres humanos de su capacidad de decisión, especialmente en situaciones críticas.
- En el diseño de los sistemas de IA es importante garantizar que los seres humanos mantengan el control sobre las decisiones importantes.

7

BENEFICIO SOCIAL ■

- El desarrollo de la IA debe contribuir al bienestar de la sociedad, promoviendo beneficios colectivos y minimizando los riesgos. Esto incluye el uso de la IA para resolver problemas globales como el cambio climático, la pobreza y las desigualdades.
- Es esencial que la IA se desarrolle y se utilice para el bien común, asegurando que los beneficios de la tecnología se distribuyan de manera equitativa entre todas las partes de la sociedad.

8

IMPACTO LABORAL ■

- La automatización impulsada por la IA puede reemplazar trabajos humanos en muchos sectores, lo que plantea cuestiones éticas sobre su impacto en el empleo y la distribución económica.
- Los responsables del desarrollo de la IA deben considerar cómo mitigar estos impactos, promoviendo políticas de reentrenamiento y educación continua para los trabajadores desplazados.

¿Conoces el 'Reglamento IA' aprobado por la UE?

El nuevo Reglamento de IA de la Unión Europea, publicado en el Diario Oficial de la Unión Europea el 12 de julio de 2024, es la normativa de alcance más ambicioso aprobada sobre esta materia hasta la fecha en el mundo, y busca impulsar la excelencia y la innovación tecnológica asegurando la protección de los derechos humanos

Europa ha sido consciente desde el primer momento de los riesgos de la IA. Por ello, el 1 de agosto de 2024 entraba en vigor en la Unión Europea (UE) la **primera ley de inteligencia artificial del mundo**.

El citado Reglamento (UE) 2024/1689 del Parlamento Europeo y del Consejo, de 13 de junio de 2024, de Inteligencia Artificial ("Reglamento IA" o "RIA") nace en el contexto de la Estrategia Europea de Inteligencia Artificial de la Comisión Europea, a través de la cual se pretende convertir a la UE en una región de **referencia mundial para la IA, garantizando que esta se centre en el ser humano y sea sostenible, segura, inclusiva y fiable, y que garantice el respeto a los derechos fundamentales, la democracia, el Estado de Derecho y la sostenibilidad medioambiental**.

A grandes rasgos, el objetivo de la nueva regulación es **limitar los riesgos de la inteligencia artificial, además de potenciar sus beneficios sociales y económicos, sin que esto suponga un obstáculo para la innovación**.

Mientras que grandes potencias como China y EEUU invierten en la investigación y el desarrollo de la IA, Europa centra sus esfuerzos en la regulación para "fomentar una tecnología fiable en Europa y fuera de ella, garantizando que los sistemas de Inteligencia Artificial respeten los derechos fundamentales, la seguridad y los principios éticos, abordando los riesgos de modelos más avanzados y potentes de IA".

El 'Reglamento IA' será de plena aplicación en Europa en 2027

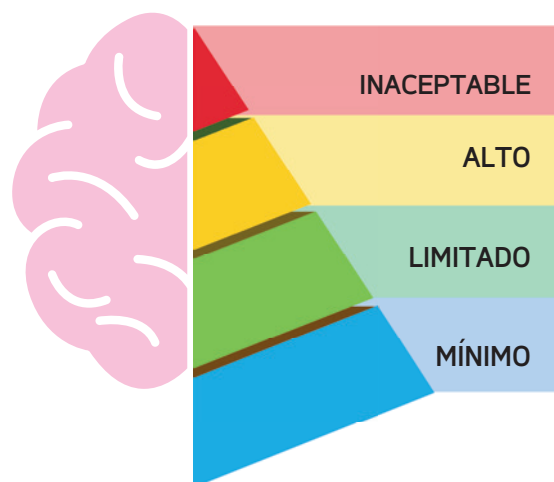
Tras su entrada en vigor el 1 de agosto de 2024, el "RIA" será de plena aplicación 24 meses después, con excepción de las prohibiciones de prácticas, que se aplicarán seis meses después de la fecha de entrada en vigor, es decir, en febrero de 2025.

En agosto de 2025, empezarán a aplicarse las normas para los modelos de uso generalista, como ChatGPT, y un año después, en agosto de 2026, se aplicará la ley de manera general, salvo algunas disposiciones.

Las obligaciones para los sistemas de alto riesgo comenzarán a aplicarse 36 meses después, en agosto de 2027.

Cuatro niveles de cumplimiento del 'Reglamento IA'

Europa pondrá multas que podrán llegar hasta los 35 millones de euros a todo el que no cumpla con la nueva normativa. Este cumplimiento se ordena en cuatro niveles de riesgo.



RIESGO MÍNIMO O NULO.

La Ley de IA permite el uso libre de la IA de riesgo mínimo. Esto incluye aplicaciones como videojuegos habilitados para IA o filtros de spam, por lo que no estarán restringidos. La gran mayoría de los sistemas de IA utilizados actualmente en la UE entran en esta categoría.

RIESGO LIMITADO.

El riesgo limitado se reduce a los riesgos asociados con la falta de transparencia en el uso de la IA. La Ley de IA introduce obligaciones específicas de transparencia para garantizar que los seres humanos estén informados cuando sea necesario, fomentando la confianza. Por ejemplo, cuando se utilizan sistemas de IA como chatbots, los seres humanos deben ser conscientes de que están interactuando con una máquina para que puedan tomar una decisión informada de continuar o dar un paso atrás. Los proveedores también tendrán que asegurarse de que el contenido generado por IA sea identificable. Además, el texto generado por IA publicado con el propósito de informar al público sobre asuntos de interés público debe etiquetarse como generado artificialmente. Esto también se aplica al contenido de audio y video.

ALTO RIESGO.

Los sistemas de IA identificados como de alto riesgo incluyen la tecnología de IA utilizada en:

INFRAESTRUCTURAS CRÍTICAS (por ejemplo, el transporte), que podrían poner en peligro la vida y la salud de los ciudadanos.

FORMACIÓN EDUCATIVA O PROFESIONAL, que puede determinar el acceso a la educación y al curso profesional de la vida de una persona (por ejemplo, puntuación de los exámenes).

COMPONENTES DE SEGURIDAD DE LOS PRODUCTOS (por ejemplo, aplicación de IA en cirugía asistida por robot).

EMPLEO, GESTIÓN DE LOS TRABAJADORES Y ACCESO AL TRABAJO POR CUENTA PROPIA (por ejemplo, programas informáticos de selección de currículos para los procedimientos de contratación).

SERVICIOS PÚBLICOS Y PRIVADOS ESENCIALES (por ejemplo, calificación crediticia que deniegue a los ciudadanos la oportunidad de obtener un préstamo).

APLICACIÓN DE LA LEY QUE PUEDA INTERFERIR CON LOS DERECHOS FUNDAMENTALES DE LAS PERSONAS (por ejemplo, evaluación de la fiabilidad de las pruebas).

GESTIÓN DE LA MIGRACIÓN, EL ASILO Y EL CONTROL FRONTERIZO (por ejemplo, examen automatizado de las solicitudes de visado).

ADMINISTRACIÓN DE JUSTICIA Y PROCESOS DEMOCRÁTICOS (por ejemplo, soluciones de IA para buscar resoluciones judiciales).

Los sistemas de IA de alto riesgo estarán sujetos a obligaciones estrictas antes de que puedan comercializarse.

RIESGO INACEPTABLE.

Los sistemas de IA que vulneran derechos fundamentales y valores de la Unión Europea se prohíben, como el respeto a la dignidad humana, la democracia y el estado de derecho. Estos sistemas están prohibidos o, en el caso de la vigilancia biométrica en tiempo real por motivos de seguridad, sujetos a estrictas restricciones.

Por ejemplo, están prohibidos los sistemas de IA que manipulan el comportamiento humano para eludir la voluntad de los usuarios o que permiten la asignación de la conocida como “puntuación social” (“social scoring”) por parte de las autoridades públicas. ■



España, ¿qué pasos está dando para afianzar la IA y proteger los derechos de la ciudadanía?

La Estrategia de Inteligencia Artificial 2024, que se desplegará durante los años 2024-2025, pretende desarrollar y expandir el uso de la inteligencia artificial de forma transparente y ética, y está enfocada a facilitar su uso en los sectores público y privado.

El Consejo de Ministros aprobó el 15 de mayo de 2024 la Estrategia de Inteligencia Artificial 2024, que da continuidad y refuerza la Estrategia de Inteligencia Artificial (ENIA) publicada en 2020, en cumplimiento del Plan de Recuperación, Transformación y Resiliencia (PRTR), y la adapta a los nuevos avances tecnológicos.

Cabe destacar que el sexto y último eje estratégico de la ENIA tiene como objetivo establecer un marco ético y normativo que refuerce la protección de los derechos individuales y colectivos, para garantizar la inclusión y el bienestar social en una sociedad en la que las inteligencias artificiales ya se están convirtiendo en algo cotidiano. Para ello se han diseñado iniciativas como:

- **Plan Nacional de Protección de Colectivos Vulnerables en IA.** Tiene como finalidad minimizar la discriminación provocada por la aplicación de sistemas inteligentes que hayan sido diseñados con sesgos.
- **Sello IA.** España acomete la creación de un sello nacional para asegurarlos sistemas que lo obtienen aplican los sistemas y algoritmos

IA de forma ética, no discriminatoria y fiable. Será un sello de aplicación voluntaria para las empresas, que demostrará que son confiables a sus clientes y usuarios.

El impulso de la inteligencia artificial también está contemplado en la Agenda España Digital 2026 como un elemento clave de carácter transversal para transformar el modelo productivo e impulsar el crecimiento de la economía española.

Asimismo, la Ley Orgánica de Protección de Datos y Garantía de Derechos Digitales (LOPDGDD), aprobada en 2018, adapta el Reglamento General de Protección de Datos (RGPD) de la UE a la legislación española y es un marco crucial para las aplicaciones de IA, ya que regula:

- **El tratamiento automatizado de datos:** Cómo los sistemas de IA pueden recopilar, almacenar y procesar datos personales.
- **Derechos digitales:** Protege los derechos de los ciudadanos frente al uso de sus datos por parte de sistemas automatizados, incluyendo el derecho a no ser objeto de decisiones basadas únicamente en el tratamiento automatizado.
- **Transparencia:** Se exige que los ciudadanos sean informados sobre el uso de sus datos personales y los algoritmos utilizados.■



La 'Carta de Derechos Digitales' española, ¿qué condiciones pone a la IA?

En el contexto de la inteligencia artificial, la **Carta de Derechos Digitales** establece un marco para proteger los derechos de la ciudadanía frente a los impactos de la IA, con la finalidad de asegurar que estas tecnologías se utilicen de manera ética y que los derechos fundamentales sean respetados en un entorno digital en constante evolución.

Ello incluye garantizar la transparencia en el uso de tecnologías de IA, asegurar la privacidad de los datos personales, prevenir sesgos y discriminación, y promover la responsabilidad y la rendición de cuentas en el desarrollo y la implementación de sistemas de IA.

Esta Carta, presentada en julio de 2021, dedica su punto 5.XXV a los Derechos ante la inteligencia artificial, señalando que:

1. La inteligencia artificial deberá asegurar un enfoque centrado en la persona y su inalienable dignidad, perseguirá el bien común y asegurará cumplir con el principio de no maleficencia.

2. En el desarrollo y ciclo de vida de los sistemas de inteligencia artificial:

- Se deberá garantizar el derecho a la no discriminación cualquiera que fuera su origen, causa o naturaleza, en relación con las decisiones, uso de datos y procesos basados en inteligencia artificial.
- Se establecerán condiciones de transparencia, auditabilidad, explicabilidad, trazabilidad, supervisión humana y gobernanza. En todo caso, la información facilitada deberá ser accesible y comprensible.
- Deberán garantizarse la accesibilidad, usabilidad y fiabilidad.

3. Las personas tienen derecho a solicitar una supervisión e intervención humana y a impugnar las decisiones automatizadas tomadas por sistemas de inteligencia artificial que produzcan efectos en su esfera personal y patrimonial.

Hay que tener en cuenta que existen normas españolas que ya reconocen la revisión y supervisión humana, como: el artículo 73.11 del Real Decreto-ley 24/2021 sobre sobre derechos de autor (3) o el artículo 14.1 de la Ley Orgánica 7/2021 sobre tratamientos de datos con fines de investigación penal.■



Proyectos de IA con enfoque ético en Gran Canaria ¿Los conoces?

En Gran Canaria, importantes proyectos que integran la inteligencia artificial (IA) están contribuyendo a la transformación digital y al desarrollo de la isla en áreas como la gestión pública, el turismo, la sanidad o la investigación.

A continuación, se citan algunos ejemplos de proyectos de inteligencia artificial (IA) que han adoptado un enfoque ético para garantizar un uso responsable de la tecnología:

■ **GobLab Gran Canaria.** Este laboratorio de innovación forma parte del Plan Estratégico de Gobernanza e Innovación Pública (PEGIP) del Cabildo de Gran Canaria. Su objetivo es integrar tecnologías avanzadas, como la IA, en la administración pública, poniendo un fuerte énfasis en la ética. Los proyectos de este laboratorio están enfocados en mejorar la atención ciudadana mediante chatbots, optimizar procesos administrativos y asegurar la transparencia y protección de datos.

■ **Aceleradora de IA para el Sector Marítimo.** Esta iniciativa, impulsada por el Cabildo de Gran Canaria en colaboración con la Autoridad Portuaria de Las Palmas y la Universidad de Las Palmas de Gran Canaria (ULPGC) se centra en aplicar tecnologías avanzadas, como la inteligencia artificial y blockchain, para optimizar las operaciones en el puerto de Las Palmas. La ética en este caso se aborda mediante la mejora de la seguridad y la transparencia en la gestión de datos y operaciones, aspectos críticos en un sector tan regulado como el marítimo.

■ **Guía ‘La inteligencia artificial (IA) en el ámbito educativo’.** El Gobierno de Canarias ha lanzado una guía sobre el uso de IA en el ámbito educativo que incorpora principios éticos como la seguridad, la privacidad y la equidad. Esta guía está orientada a asegurar que la IA se utilice de manera inclusiva y responsable en las aulas, protegiendo los derechos de los estudiantes y asegurando que las herramientas tecnológicas contribuyan a mejorar el aprendizaje sin poner en riesgo los valores fundamentales.

■ **La Revista Canaria de Administración Pública** ha publicado investigaciones sobre las cuestiones éticas que surgen con la adopción de la IA en las instituciones gubernamentales. Estas discusiones abarcan aspectos como la privacidad, la seguridad, la explicabilidad y la responsabilidad en la toma de decisiones automatizadas. También se analizan los retos que enfrenta el sector público en cuanto a la implementación ética de la IA.

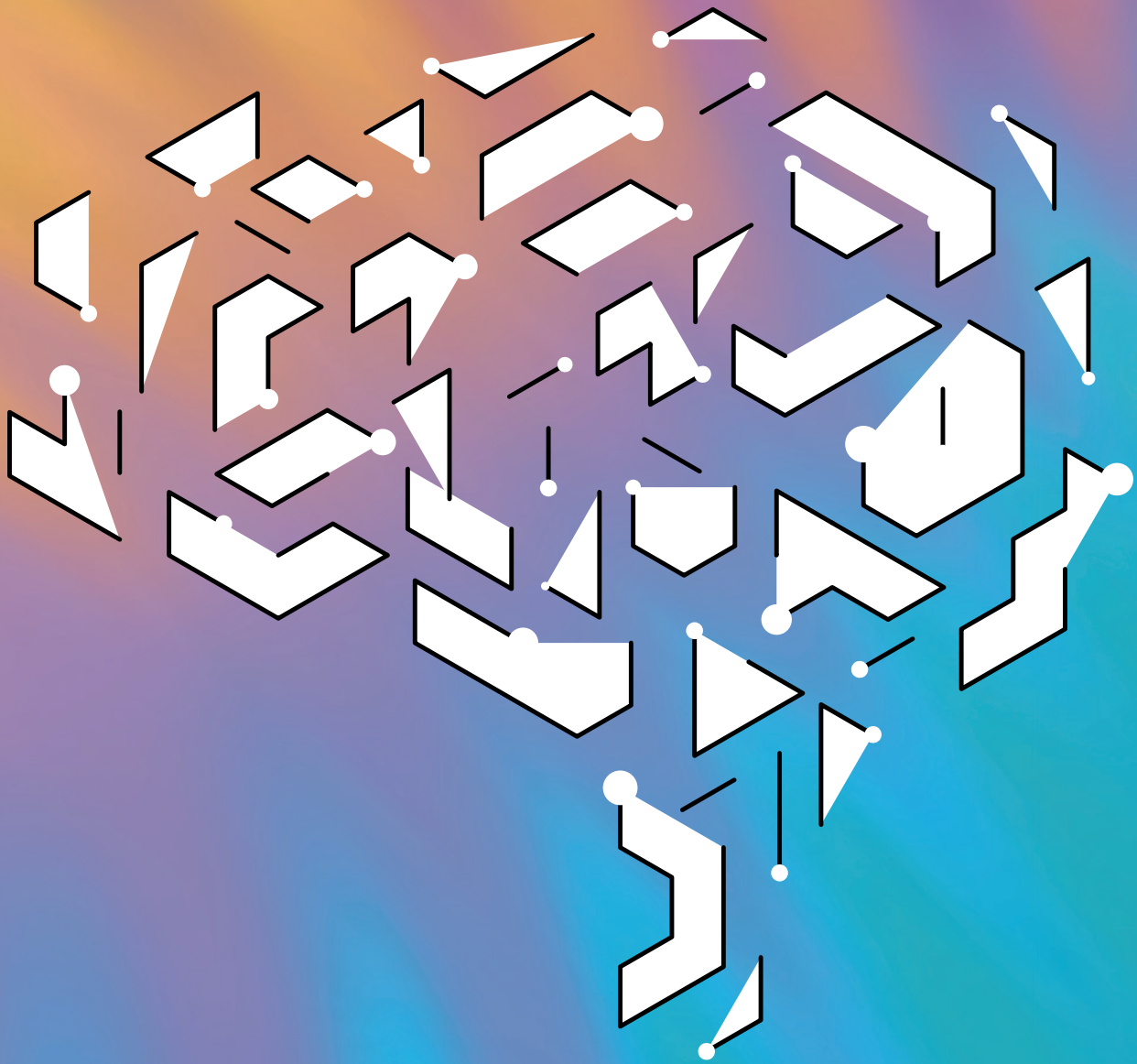
■ **Proyecto ‘Inteligencia Artificial: Información, Sensibilidad y Ética para una Ciudadanía Consciente’.** Desarrollado por la Asociación Gran Canaria para la UNESCO, este proyecto pretende impulsar, “de forma más urgente que nunca”, la implementación a nivel local, concretamente en Gran Canaria, de la IA basada en principios éticos, de acuerdo al marco general que alienta la UNESCO. Ante los riesgos que los sistemas de IA representan para la sociedad, parte de ella exacerbada por la ausencia de sensibilización e información clara, se antoja imprescindible la creación de contenido que ponga luz y permita conocer cómo debe operar la Inteligencia Artificial, y su plasmación en un boletín digital anual, titulado ‘Panorama IA’, para su difusión entre la sociedad civil de Gran Canaria.



En ninguna otra especialidad necesitamos más una “brújula ética” que en la inteligencia artificial. Estas tecnologías de utilidad general están remodelando nuestra forma de trabajar, interactuar y vivir. El mundo está a punto de cambiar a un ritmo que no se veía desde el despliegue de la imprenta hace más de seis siglos. La tecnología de inteligencia artificial aporta grandes beneficios en muchos ámbitos, pero sin unas barreras éticas corre el riesgo de reproducir los prejuicios y la discriminación del mundo real, alimentar las divisiones y amenazar los derechos humanos y las libertades fundamentales.■

Gabriela Ramos,
Subdirectora General de Ciencias Sociales y
Humanas de la UNESCO





FUENTES CONSULTADAS:

CIECEM

Comisión Europea – Ley de IA

ChatGPT

GobLab Gran Canaria

GulA de buenas prácticas en el uso de la inteligencia artificial ética, OdiselA.

Observatorio del Impacto Social y Ético de la Inteligencia Artificial

Ministerio para la Transformación Digital y Función Pública – Gobierno de España

Puerto Canarias

Revista Canaria de Administración Pública

UNESCO