
INTELIGENCIA ARTIFICIAL

Los 7 mandamientos de la inteligencia artificial que rigen su futuro

Por Netvez · 26 de julio de 2026 · 10 min de lectura

Descubre los 7 mandamientos de la IA que Google, Microsoft y la UNESCO aplican en silencio. Principios éticos verificables que cambian cómo entiendes la inteligencia artificial.

La inteligencia artificial ya no es una promesa de laboratorio; está en tu teléfono, en tu correo, en la decisión de si recibes un préstamo y en el diagnóstico que lee un médico. Con tanta presencia, surge una pregunta incómoda: ¿quién le puso reglas a la máquina? Los 7 mandamientos de la IA no son un decálogo religioso ni un manual de programación. Son principios éticos que organismos como la UNESCO, la Unión Europea y empresas como Google o Microsoft han asumido como base de su desarrollo, aunque pocos usuarios los conozcan. Entender estos mandamientos no es tarea solo de ingenieros. Es una necesidad para cualquier profesional, estudiante o lector que quiera saber qué límites invisibles rigen la tecnología que ya forma parte de su día a día.

Por qué la [inteligencia artificial](#) necesita mandamientos en lugar de simples reglas

La tecnología funciona con código, pero la sociedad funciona con valores. Cuando una inteligencia artificial decide quién entra a la universidad, quién recibe una beca o qué noticia es relevante, no basta con que el algoritmo sea eficiente; debe ser justo. Esa distinción separa una regla técnica de un mandamiento ético. Las reglas le dicen a la máquina cómo hacer algo; los mandamientos le dicen qué no debe hacer nunca. En 2019, el investigador Juan Antonio Lloret Egea publicó un paper titulado exactamente Los 7 mandamientos de la IA, donde

recogía principios que iban desde la auditoría universal de los datos hasta la necesidad de un marco legal global. Esos puntos no eran fantasía. Hoy, la Ley de Inteligencia Artificial de la Unión Europea y las recomendaciones de la UNESCO los han convertido en norma. Entender este origen es clave para leer el resto del artículo con la perspectiva adecuada: no estamos ante opiniones personales, sino ante directrices que ya mueven millones de dólares y afectan a miles de millones de personas.

El legado de Asimov y el nacimiento de una ética real para la IA (inteligencia artificial)

En 1942, Isaac Asimov escribió las tres leyes de la robótica. Un robot no debe dañar a un ser humano, debe obedecer las órdenes y proteger su propia existencia siempre que no entre en conflicto con las dos primeras. Durante décadas, esas leyes fueron el único referente popular sobre ética en máquinas, pero eran ficción. La realidad exige algo más complejo: [no basta con que un algoritmo](#) no dañe de forma directa; debe evitar dañar de forma indirecta, sesgada o invisible. Los principios de la IA que no todos conocen nacieron de esa necesidad. No son leyes de un relato de ciencia ficción, sino pautas que organismos internacionales redactan en documentos técnicos y que las empresas tecnológicas incorporan en sus políticas internas. Asimov soñó con robots obedientes; hoy necesitamos algoritmos responsables.

El paper (investigación) que recogió los siete principios antes de que la ley europea los hiciera obligatorios para los creadores de inteligencia artificial

El trabajo de Lloret Egea, publicado en engrxiv.org en 2019, anticipó debates que hoy parecen actuales. Su propuesta incluía que los estándares de la inteligencia artificial debieran ser abiertos y auditables, que la regulación debiera ser universal con la ONU como eje central, y que los datos manejados por la IA debieran someterse a auditoría mundial para evitar sesgos territoriales. Lo llamativo es que esos puntos no eran utopías técnicas. En 2021, la UNESCO adoptó su Recomendación sobre la Ética de la Inteligencia Artificial, donde aparecen ideas paralelas sobre transparencia, justicia y gobernanza. La Unión Europea, por su parte, aprobó en 2024 su Ley de Inteligencia Artificial, que clasifica los sistemas por nivel de riesgo y exige

supervisión humana en los de alto impacto. El paper de 2019 ya apuntaba hacia ese horizonte regulatorio.

Los tres mandamientos que las grandes tecnológicas asumen pero no explican

Cuando usas un buscador, un asistente de voz o una herramienta de traducción, estás interactuando con decisiones que ya han sido filtradas por principios éticos internos. Google, Microsoft, IBM y otras grandes corporaciones publican sus directrices de inteligencia artificial responsable, aunque rara vez las explican al usuario final. Esos documentos no son marketing; son compromisos que afectan qué proyectos se financian, qué datos se usan y qué productos llegan al mercado. Conocer estos tres primeros mandamientos de la IA es entender qué valores se programan en silencio cada vez que pulsas enter.

Primero: la IA debe rendir cuentas ante las personas, no ante sí misma

Google estableció en 2018 que sus sistemas de inteligencia artificial deben rendir cuentas a las personas. Microsoft usa el término responsabilidad algorítmica. La idea es la misma: si un algoritmo rechaza tu solicitud de empleo, niega un crédito o marca un contenido como spam, debe existir un mecanismo para que un humano revise esa decisión. No se trata de que la máquina se disculpe, sino de que su funcionamiento pueda ser explicado, auditado y, en caso de error, corregido. La rendición de cuentas es el primer mandamiento porque sin él los demás no tienen sentido. Un algoritmo opaco no puede ser justo, aunque sus resultados parezcan correctos.

Segundo: la IA debe evitar crear o reforzar sesgos injustos

Los sesgos en la inteligencia artificial no son errores técnicos; son reflejos de datos históricos que la máquina aprende y amplifica. Si un sistema de contratación se entrena con currículos de una empresa donde históricamente predominan los hombres, tenderá a penalizar currículos femeninos. Si un sistema de justicia predictiva usa código postal como variable, puede

discriminar indirectamente por raza o clase social. Google, Microsoft y la Unión Europea coinciden en que la inteligencia artificial debe ser diseñada para detectar, mitigar y prevenir estos sesgos. No basta con ser neutral; debe ser activamente equitativa. Este segundo mandamiento exige que los equipos de desarrollo revisen sus datos, diversifiquen sus equipos y prueben sus modelos contra escenarios de discriminación antes de lanzarlos.

Tercero: la IA debe preservar la privacidad por diseño, no como añadido

La privacidad no puede ser una casilla que se marca al final del desarrollo; debe estar en el código desde el primer día. Microsoft llama a esto privacidad inteligente; la Unión Europea lo denomina privacidad por diseño y por defecto. Significa que un sistema de inteligencia artificial debe recolectar solo los datos estrictamente necesarios, minimizar la información personal y garantizar que el usuario controla qué se comparte y con quién. Cuando una IA generativa guarda tus conversaciones para entrenar futuros modelos, o cuando un asistente de voz escucha más allá de la palabra de activación, está violando este mandamiento. Las normas de la inteligencia artificial exigen que la privacidad sea inherente, no opcional.

Lo que Europa y la UNESCO convirtieron en norma internacional para la inteligencia artificial

Mientras las empresas tecnológicas redactaban sus principios internos, los gobiernos y organismos multilaterales hacían lo propio a escala global. La UNESCO, con sus 193 Estados miembros, adoptó en 2021 la primera norma universal sobre ética de la inteligencia artificial. La Unión Europea aprobó en 2024 la primera ley integral del mundo sobre la materia. Estos textos no son sugerencias; son marcos que condicionan inversión, innovación y comercio. Los tres mandamientos que vienen a continuación no son deseos de un manifiesto; son obligaciones legales que ya rigen cómo se construye y se despliega la inteligencia artificial en buena parte del planeta.

Cuarto: la IA debe ser transparente en su funcionamiento

La transparencia en la inteligencia artificial significa que los desarrolladores deben publicar información sobre el uso, las limitaciones y el posible impacto de sus sistemas. La UNESCO lo incluye como pilar fundamental; la Comisión Europea lo exige para los sistemas de alto riesgo. Un usuario tiene derecho a saber cuándo está interactuando con una IA, no con un humano. Un afectado por una decisión algorítmica tiene derecho a una explicación comprensible. La transparencia no implica revelar cada línea de código, pero sí garantizar que el proceso de decisión pueda ser interpretado por personas sin formación técnica. Sin este cuarto mandamiento, la inteligencia artificial opera como una caja negra, y las cajas negras no pueden ser democráticas.

Quinto: la IA debe ser socialmente beneficiosa, no solo eficiente

Google abre sus principios de inteligencia artificial con esta frase: la IA debe ser socialmente beneficiosa. Parece obvio, pero no lo es. Un algoritmo puede ser extremadamente eficiente en predecir qué producto comprarás, qué vídeo te enganchará o cuánto pagarás por un seguro, y aun así ser socialmente perjudicial. La eficiencia sin propósito ético ha llevado a sistemas de recomendación que radicalizan opiniones, a modelos de precios dinámicos que penalizan la pobreza y a herramientas de vigilancia que erosionan libertades. El quinto mandamiento exige que cada sistema de inteligencia artificial se evalúe no solo por su rendimiento técnico, sino por su impacto social. Debe crear oportunidades, reducir desigualdades y respetar los derechos humanos.

Sexto: la IA debe respetar la autonomía humana en cada decisión

La autonomía humana es el principio de que la persona debe estar siempre en el centro. No se trata de que la IA sea una herramienta pasiva, sino de que nunca sustituya la capacidad de decisión de un ser humano en contextos que le afecten directamente. La Unión Europea lo denomina supervisión humana; Microsoft lo enmarca dentro de la ayuda a la humanidad respetando su autonomía. Un médico puede usar una IA para interpretar una radiografía, pero la decisión final debe ser suya. Un juez puede consultar un algoritmo predictivo, pero la sentencia debe llevar su firma. La inteligencia artificial que respeta este mandamiento no

reemplaza el juicio humano; lo amplifica. Cuando un sistema automatiza decisiones críticas sin posibilidad de apelación, viola este principio.

El séptimo mandamiento y la frontera que aún no hemos cruzado

Los seis mandamientos anteriores ya aparecen en documentos corporativos, leyes y recomendaciones internacionales. El séptimo es más incómodo porque apunta a un territorio que la regulación aún no ha cubierto del todo. Habla de la necesidad de que la inteligencia artificial sea auditada de forma universal, que sus estándares sean abiertos y que exista un organismo internacional capaz de certificar tanto los productos como las personas que los desarrollan. No es ciencia ficción. Es la propuesta que Lloret Egea ya formuló en 2019 y que hoy resuena en los debates sobre gobernanza global de la tecnología.

Séptimo: la IA no debe delegar decisiones críticas sin supervisión humana

Este mandamiento no se limita a repetir que un humano debe estar en el circuito; va más allá. Exige que exista un marco legal universal, coordinado por organismos como la ONU, que garantice que ningún sistema de inteligencia artificial tome decisiones de alto impacto sin que un humano pueda intervenir, corregir o anular el resultado. La Ley de Inteligencia Artificial de la Unión Europea ya exige esto para sistemas de alto riesgo, como los que evalúan exámenes o deciden sobre préstamos. Pero el reto está en que esa supervisión sea real, no simbólica. Un botón de apelación que nadie sabe usar no es supervisión. La inteligencia artificial del futuro necesitará no solo principios, sino mecanismos concretos que aseguren que el ser humano siga siendo el último responsable.

Cómo estos siete principios cambian la forma en que usas la inteligencia artificial para siempre

Conocer los 7 mandamientos de la IA [transforma la experiencia](#) del usuario. Cuando entiendes que existe un principio de transparencia, exiges saber por qué una plataforma te muestra cierto contenido. Cuando sabes que la privacidad debe ser por diseño, rechazas servicios que

piden más datos de los necesarios. Cuando internalizas que la autonomía humana es inviolable, no aceptas decisiones algorítmicas sin posibilidad de reclamación. Estos principios no son abstractos; son herramientas de empoderamiento. La inteligencia artificial seguirá evolucionando, pero su dirección depende de que los usuarios conozcan las reglas que ya existen y exijan que se cumplan. El futuro de la tecnología no lo escriben solo los programadores; lo escriben también quienes saben qué preguntar.