



2.

**IA Simbólica y generativa
LLM**

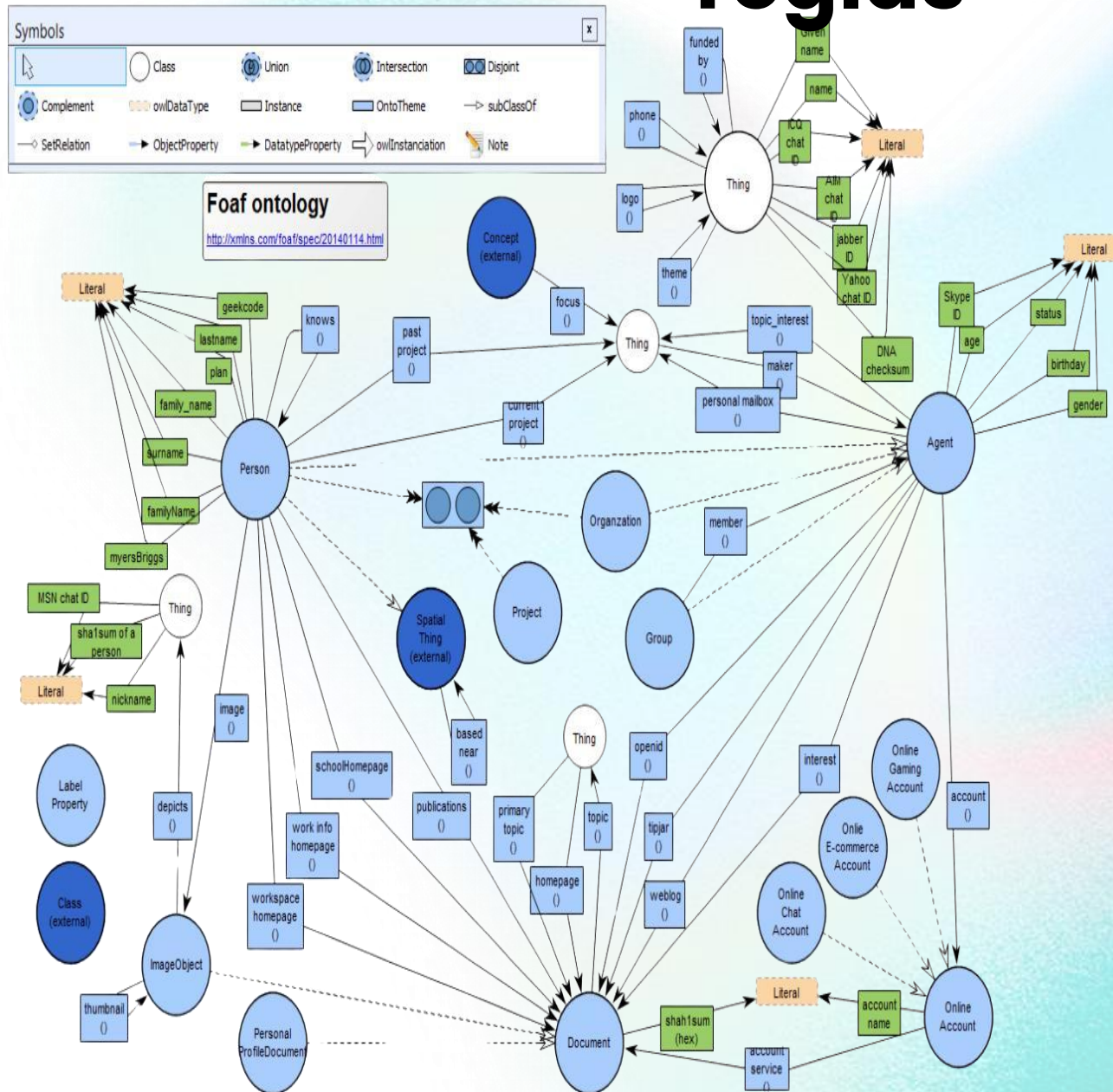
Marcos normativos IA



Modelos



Modelos de IA Simbólica basados en reglas



- El conocimiento se modela en un formato amigable computacionalmente hablando.
- Se identifican conceptos y relaciones entre conceptos.
- Se deriva en objetos y relaciones entre objetos. (Sistema)
- Los lenguajes de programación modernos también siguen esta aproximación. (Coherencia)

Algoritmo

Un **algoritmo en Inteligencia Artificial** es un **conjunto ordenado de instrucciones o reglas** que una máquina sigue para **resolver un problema, tomar decisiones o aprender a partir de datos**.

En este contexto, no es solo una secuencia fija de pasos, sino que muchas veces incorpora **capacidad de aprendizaje, adaptación y mejora continua**.

En términos simples: es el “método” que usa la IA para pasar de datos de entrada → a resultados útiles (predicciones, clasificaciones, decisiones)

Ejemplo:

- Un algoritmo puede analizar miles de correos y decidir si son **spam o no spam**.
- Otro puede aprender a reconocer **rostros en imágenes**.
- Otro puede predecir **accidentes viales** (alineado con tu línea de investigación).

Un algoritmo en IA no solo ejecuta instrucciones....**construye conocimiento a partir de la experiencia (datos)**.

Representa el **funcionamiento del cerebro humano como un sistema de procesamiento de información**, y es muy útil para explicar cómo funcionan los **algoritmos en Inteligencia Artificial**, porque muestra el modelo que la IA intenta imitar.

¿CÓMO FUNCIONAN LOS ALGORITMOS EN IA?

El cerebro humano como modelo de procesamiento de información

Un algoritmo en IA imita al cerebro: recibe datos, los **procesa**, **toma decisiones** y **aprende** de la experiencia.

1 ENTRADA DE DATOS (INPUTS)



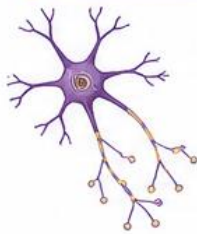
El cerebro recibe información a través de los 8 sentidos / entradas sensoriales.

En IA equivale a:

- Datos de entrada
- Ej: imágenes, texto, sensores, bases de datos

IDEA CLAVE: Sin datos, no hay algoritmo.

2 NEURONAS Y CONEXIONES



Miles de millones de neuronas conectadas entre sí.

En IA se traduce en:

- Redes neuronales artificiales
- Nodos conectados que procesan información

IDEA CLAVE: El algoritmo no trabaja aislado, sino como una red de relaciones.

3 PROCESAMIENTO (EL "ALGORITMO")



El cerebro analiza, compara, integra y transforma la información.

En IA es el algoritmo en acción:

- Aplica reglas
- Ajusta parámetros
- Calcula probabilidades

Ejemplo: Reconocer una cara → comparar patrones → decidir si coincide

5 APRENDIZAJE (FEEDBACK)



El cerebro aprende mediante experiencia, error y repetición.

En IA:

- Ajuste de pesos
- Entrenamiento del modelo
- Mejora continua

Ejemplo: Si se equivoca → corrige → mejora el resultado futuro

4 TOMA DE DECISIONES (OUTPUT)



El cerebro decide y genera una respuesta:

- Recompensa
- Castigo
- Emoción
- Acción

En IA:

- Clasificación (spam / no spam)
- Predicción (riesgo, accidentes)
- Recomendación (qué mostrar)

IDEA CLAVE: El algoritmo siempre produce una salida (output).

6 REPRESENTACIÓN DEL CONOCIMIENTO



El cerebro construye conceptos, relaciones y mapas mentales.

En IA:

- Modelos
- Variables
- Features (características)

Conecta con: "conceptos y relaciones entre conceptos".

7 VISIÓN SISTÉMICA (SISTEMA COMPLEJO)



El cerebro no trabaja por partes aisladas, funciona como un sistema integrado.

En IA:

- Un algoritmo es parte de un sistema mayor
- Interactúa con datos, contexto y objetivos

IDEA CLAVE: La IA no solo calcula, interactúa y se adapta.

APLICACIÓN REAL



Prevención de accidentes viales



Diagnóstico médico



Detección de spam



Recomendaciones personalizadas



SÍNTESIS: Un algoritmo en IA funciona como una versión simplificada del cerebro: recibe datos, los procesa mediante conexiones, toma decisiones y aprende de la experiencia.



¿CÓMO FUNCIONAN LOS ALGORITMOS EN INTELIGENCIA ARTIFICIAL?



Un algoritmo en IA recibe datos, los procesa, toma decisiones y aprende para mejorar sus resultados.

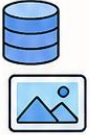
1 ENTRADA DE DATOS

El algoritmo recibe información del mundo real a través de los sentidos o dispositivos.



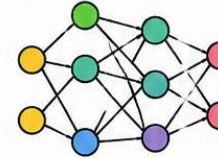
Ejemplo:

Imágenes, sonidos, texto, sensores, bases de datos.



2 PROCESAMIENTO

El algoritmo analiza los datos usando reglas y conexiones (similar a las neuronas del cerebro).



Ejemplo:

La IA detecta patrones en los datos y los transforma en información útil.



3 TOMA DE DECISIONES

Con base en lo que aprendió, el algoritmo toma una decisión o hace una predicción.



Ejemplo:

Clasifica un correo como spam o no spam. Predice el riesgo de accidente.



4 SALIDA (OUTPUT)

El algoritmo entrega un resultado que puede ser una respuesta, acción o recomendación.



Ejemplo:

Mostrar un resultado, enviar una alerta, recomendar una ruta.



5 FEEDBACK (RETROALIMENTACIÓN)

El sistema recibe información sobre si la decisión fue correcta o no.



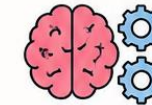
Ejemplo:

El usuario indica si la recomendación fue útil o no.



6 APRENDIZAJE

El algoritmo ajusta sus reglas internas para mejorar en el futuro.



Ejemplo:

La IA aprende de sus errores y mejora sus predicciones.



7 MEJORA CONTINUA

Con más datos y experiencia, el algoritmo se vuelve más preciso y confiable.



Ejemplo:

Con el tiempo, detecta mejor los riesgos y previene más accidentes.



EN RESUMEN

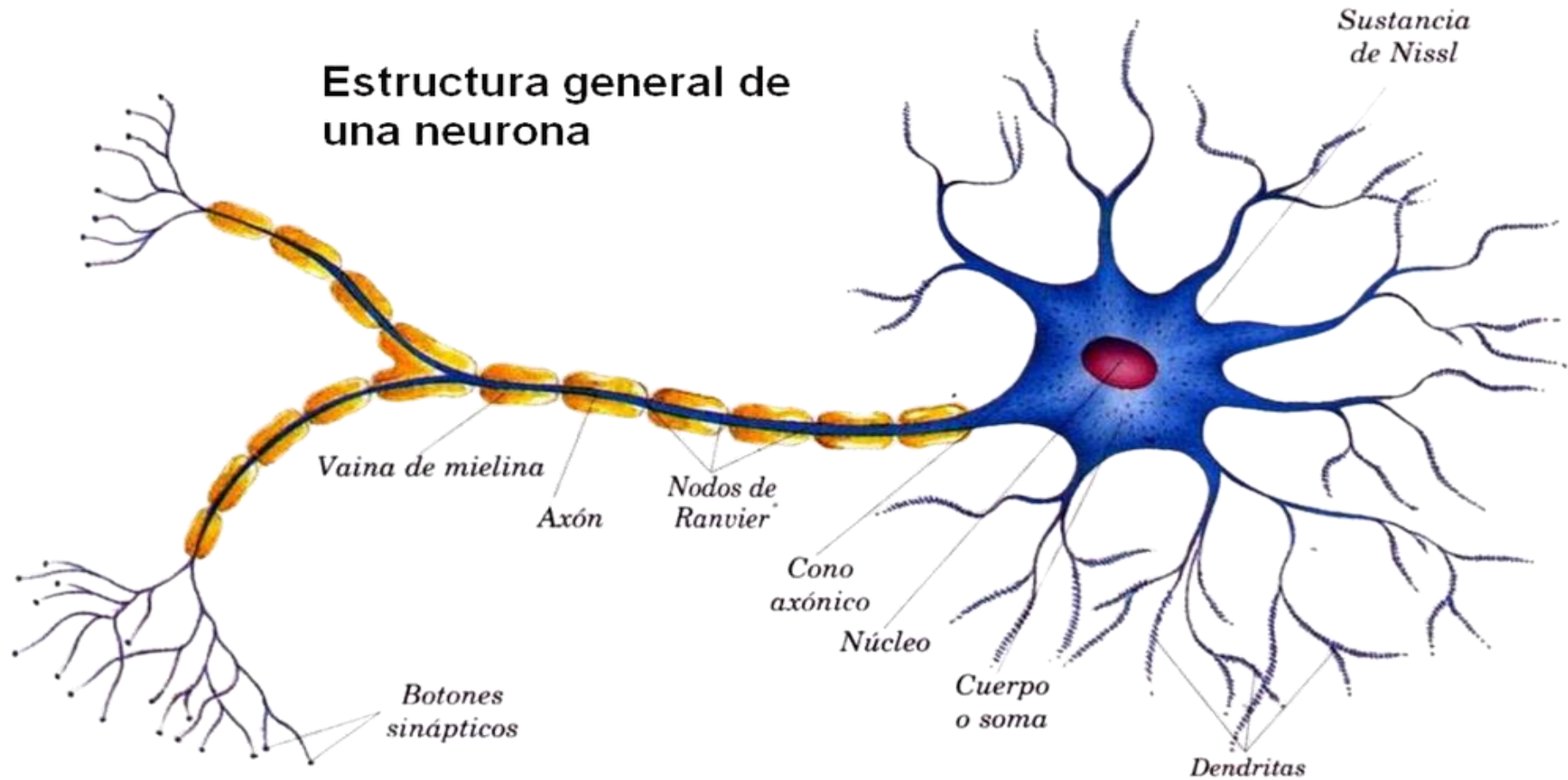
Un algoritmo en IA funciona como un ciclo:

Datos → Procesamiento → Decisión → Resultado → Aprendizaje → Mejora

Así, la inteligencia artificial aprende como los humanos: con experiencia.

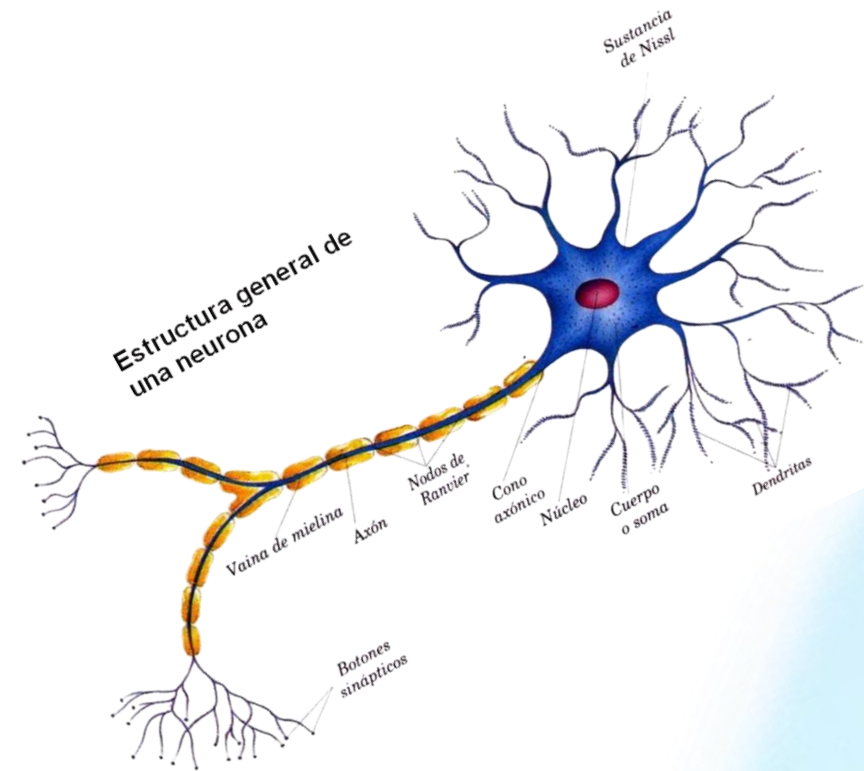


Estructura general de una neurona

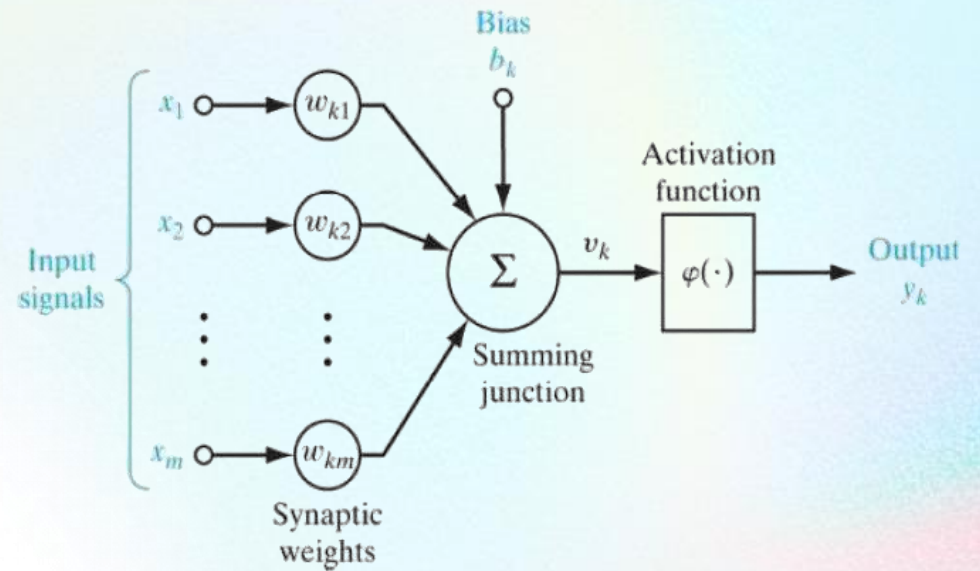
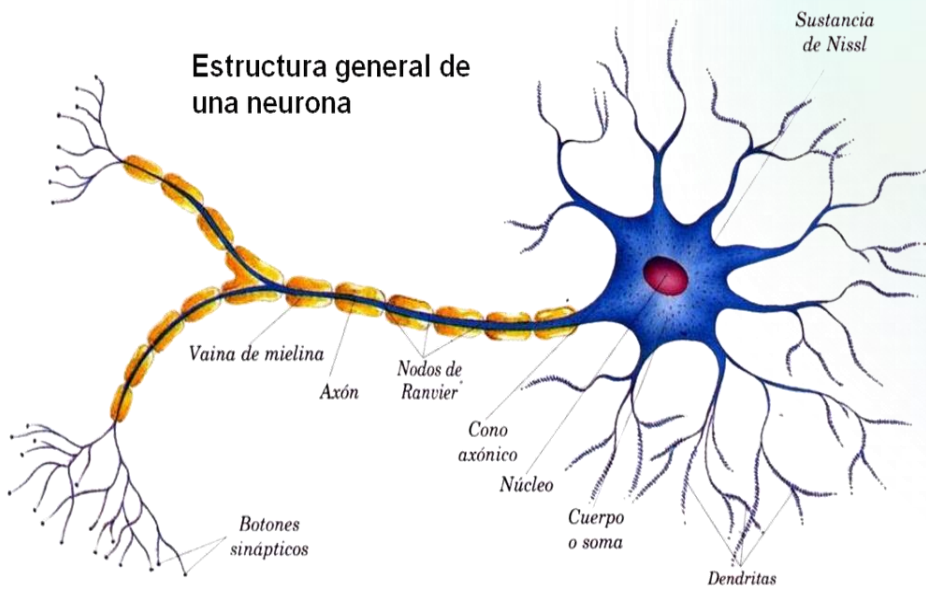


Estructura de una neurona, que es la unidad básica del cerebro. Es clave para entender la IA porque las redes neuronales artificiales están inspiradas directamente en este modelo.

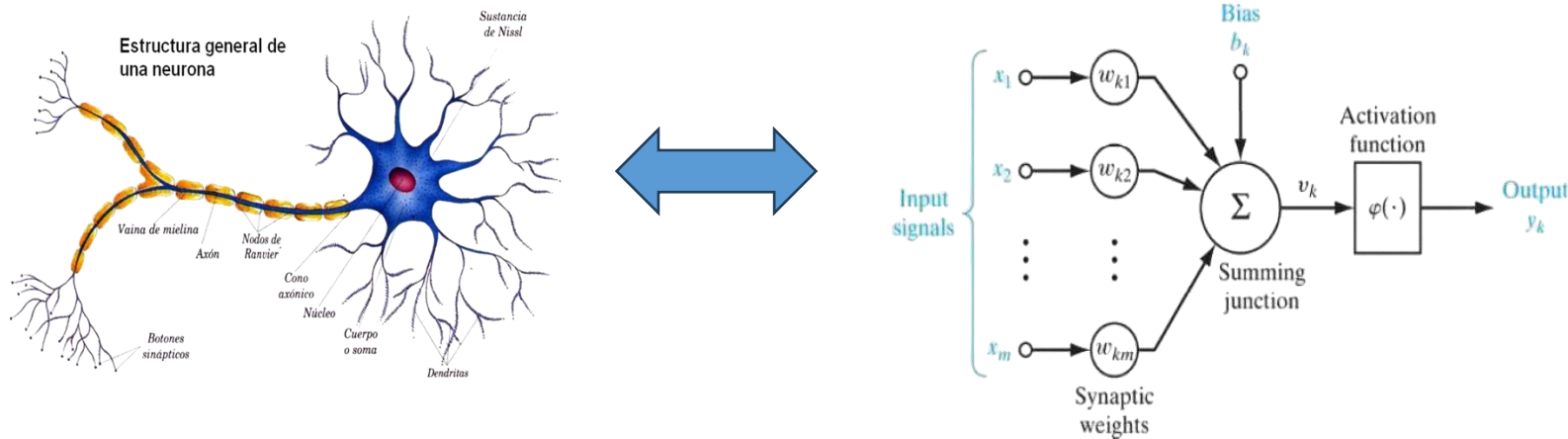
Nuestro cerebro y las neuronas



Analogía biológica y matemática: neuronas

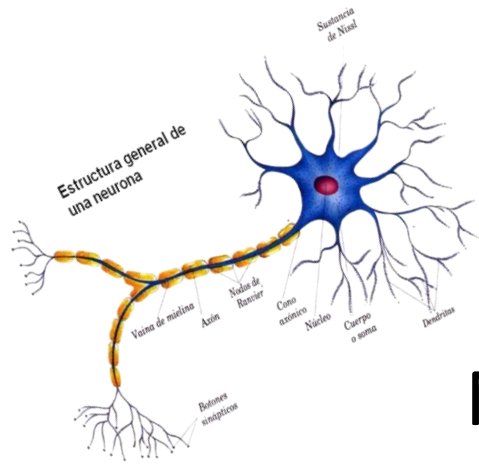


1. **Dendritas:** entrada de la información. Inputs del sistema.
2. **Cuerpo o soma:** procesamiento. Algoritmo que procesa.
3. **Núcleo:** control del funcionamiento de la neurona. Es la lógica del modelo
4. **Axón:** es la transmisión/canal hacia otra neurona. Output del procesamiento
5. **Vaina de mielina:** Acelerador de transmisión de la señal. Eficiencia y optimización del algoritmo.
6. **Nodos de Ranvier:** la señal viaja mas rápido en saltos. Optimizaciones o mejoras en flujos de datos
7. **Botones simpáticos:** conexión entre neuronas.



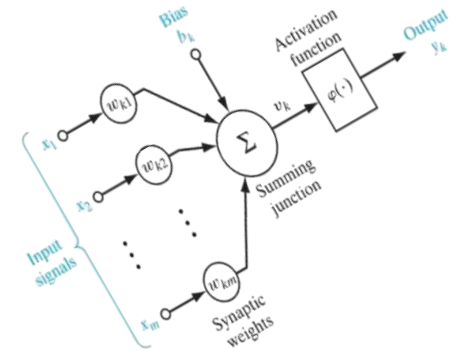
Relación directa con la IA

Una **neurona artificial:** recibe datos $\rightarrow$ los procesa $\rightarrow$ decide si “activar” o no $\rightarrow$ envía resultado
Esto es prácticamente lo mismo que hace la neurona biológica.



Neurona biológica → inspiración
 Neurona artificial → implementación
 Red neuronal → sistema inteligente

La IA no copia el cerebro exactamente, pero sí toma su lógica de funcionamiento.

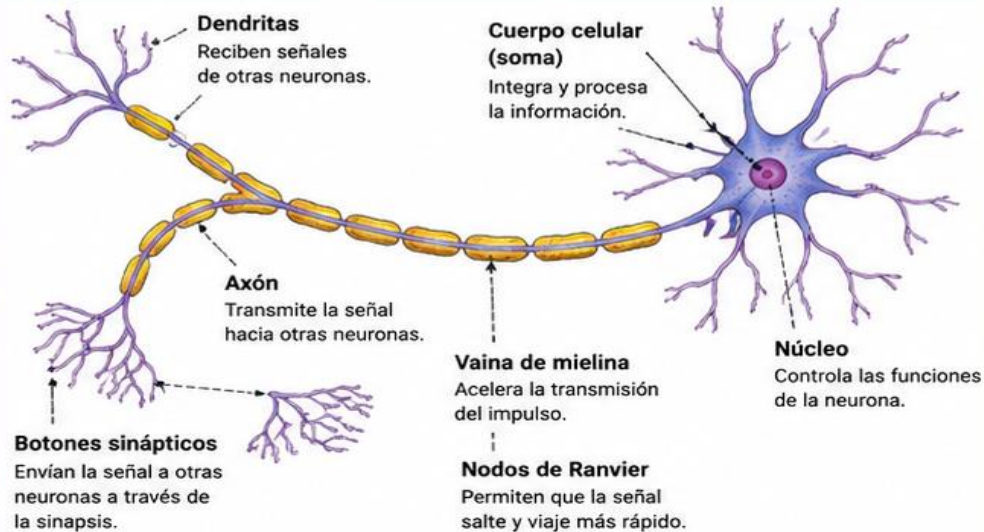


NEURONA BIOLÓGICA vs NEURONA ARTIFICIAL

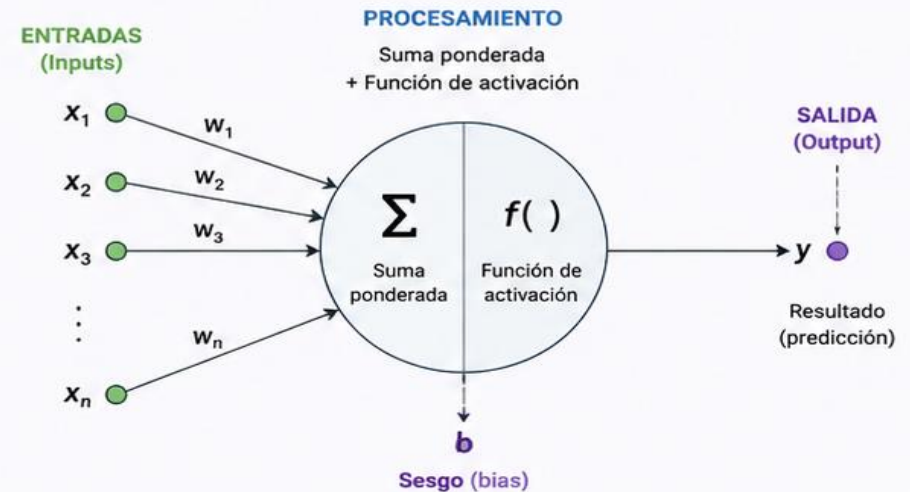


Las redes neuronales artificiales están inspiradas en cómo funciona nuestro cerebro.

NEURONA BIOLÓGICA



NEURONA ARTIFICIAL



COMPARACIÓN

NEURONA BIOLÓGICA

EQUIVALENTE EN IA

NEURONA ARTIFICIAL



Dendritas: reciben señales.

Entrada de información

Entradas ($x_1, x_2, \dots, x_n$): reciben datos.



Cuerpo celular (soma): integra y procesa.

Procesamiento

Suma ponderada + función de activación: procesa los datos.



Núcleo: controla funciones.

Parámetros del modelo

Sesgo (bias): ajusta el comportamiento de la neurona.



Axón: transmite la señal.

Salida del procesamiento

Salida (y): resultado o predicción.



Vaina de mielina: acelera la transmisión.

Optimización y eficiencia

Mejores algoritmos, técnicas y hardware aceleran el aprendizaje.



Nodos de Ranvier: permiten saltos rápidos.

Mejora del flujo de información

Optimizaciones que permiten entrenar más rápido y con mejor rendimiento.



Botones sinápticos: conectan con otras neuronas.

Conexiones y pesos

Pesos ($w_1, w_2, \dots, w_n$): determinan la fuerza de cada conexión.



IDEA CLAVE

Las redes neuronales artificiales imitan el principio de funcionamiento del cerebro:



Reciben datos



Los procesan



Toman decisiones



Aprenden y mejoran



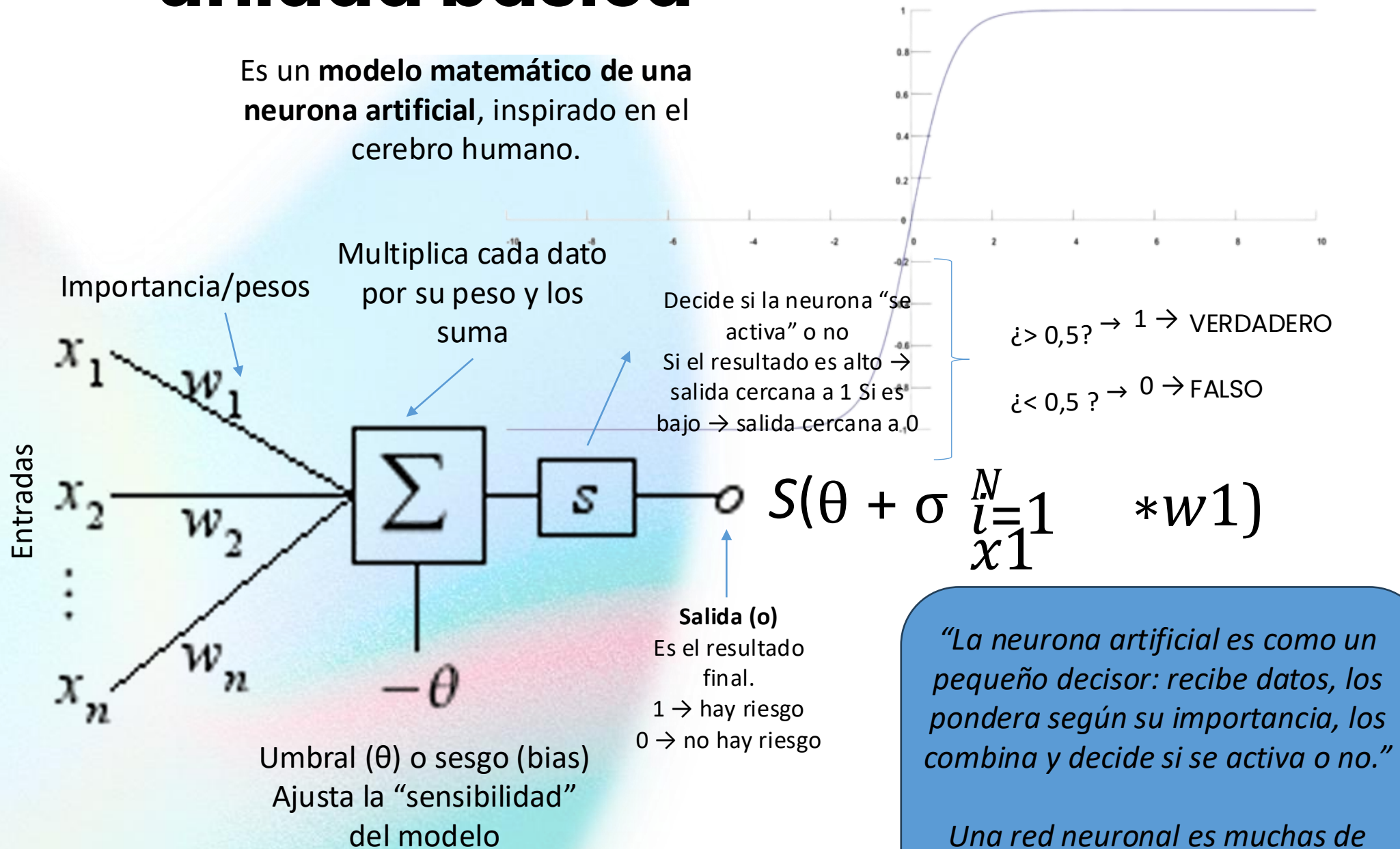
EN RESUMEN:

Una red neuronal artificial es un conjunto de neuronas artificiales conectadas entre sí que trabajan juntas para resolver problemas, aprender de los datos y tomar decisiones.



La neurona como unidad básica

Es un **modelo matemático de una neurona artificial**, inspirado en el cerebro humano.



"La neurona artificial es como un pequeño decisor: recibe datos, los pondera según su importancia, los combina y decide si se activa o no."

Una red neuronal es muchas de estas neuronas trabajando juntas.



Ejercicio práctico: neurona artificial simple



Una neurona artificial debe decidir si existe **riesgo de accidente vial**.

1. DATOS DE ENTRADA

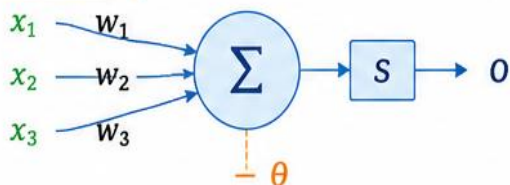
	Valor	Peso
 $x_1 =$ velocidad alta	$x_1 = 1$	$w_1 = 0,5$
 $x_2 =$ lluvia	$x_2 = 1$	$w_2 = 0,3$
 $x_3 =$ poca iluminación	$x_3 = 0$	$w_3 = 0,2$

2. UMBRAL (θ)

$\theta = 0,6$

Define el punto a partir del cual la neurona se activa.

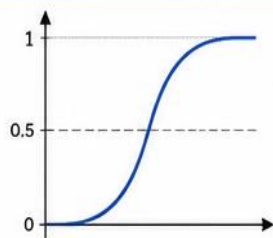
3. FÓRMULA



La neurona calcula la suma ponderada y luego aplica la función de activación $S()$.

$$\text{Resultado} = (x_1 \times w_1) + (x_2 \times w_2) + (x_3 \times w_3)$$

4. FUNCIÓN DE ACTIVACIÓN (S)



Si el resultado es alto $\rightarrow$ salida cercana a 1
Si el resultado es bajo $\rightarrow$ salida cercana a 0

Si Resultado $\geq \theta \rightarrow$ salida = 1 $\rightarrow$ Hay riesgo
Si Resultado $< \theta \rightarrow$ salida = 0 $\rightarrow$ No hay riesgo

PARA RESOLVER

1 Cálculo:

$$\text{Resultado} = (x_1 \times w_1) + (x_2 \times w_2) + (x_3 \times w_3)$$

$$\text{Resultado} = (1 \times 0,5) + (1 \times 0,3) + (0 \times 0,2)$$

$$\text{Resultado} = \underline{\hspace{2cm}}$$

2 Comparación con el umbral ($\theta = 0,6$):

¿La neurona se activa?

3 Salida:

$$\text{Salida (o)} = \underline{\hspace{2cm}} \quad (1 = \text{Hay riesgo}, 0 = \text{No hay riesgo})$$

4 Conclusión:

Hay riesgo de accidente vial.

RESPUESTA

1 Cálculo:

$$\text{Resultado} = (1 \times 0,5) + (1 \times 0,3) + (0 \times 0,2)$$

$$\text{Resultado} = 0,5 + 0,3 + 0 \rightarrow \text{Resultado} = 0,8$$

2 Comparación con el umbral:

$$0,8 \geq 0,6 \rightarrow \text{La neurona se activa.}$$

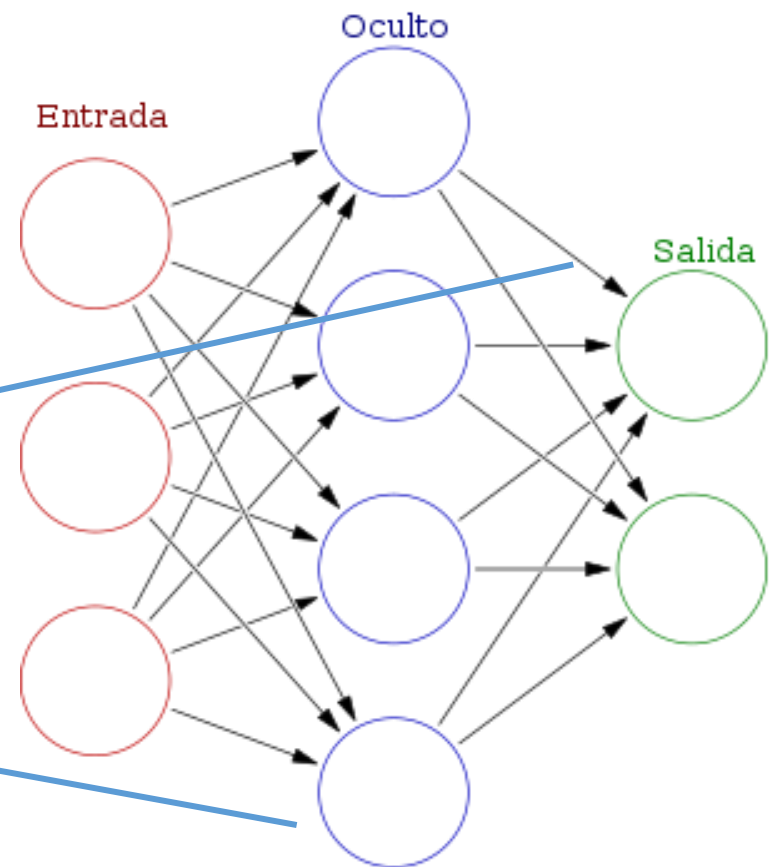
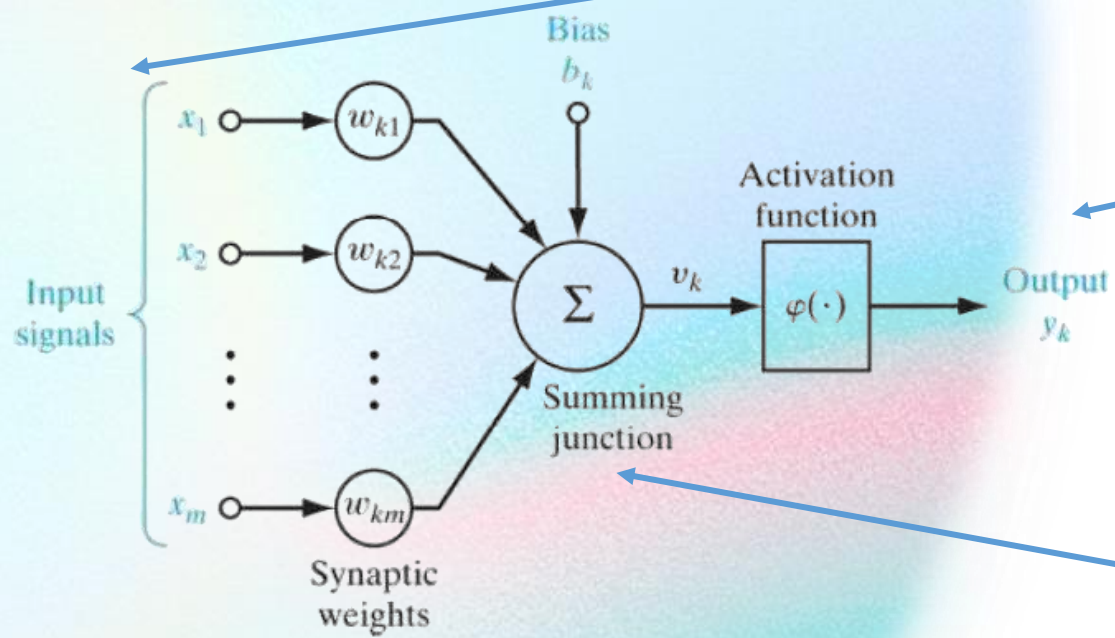
3 Salida:

$$\text{Salida (o)} = \underline{1} \quad (1 = \text{Hay riesgo}, 0 = \text{No hay riesgo})$$

4 Conclusión:

Hay riesgo de accidente vial.





Avanzando hacia la Inteligencia Artificial

Inteligencia Artificial

La ciencia y la ingeniería de fabricar máquinas inteligentes.

IA es la forma más amplia de simular la inteligencia humana en las máquinas, incluyendo la planificación, el aprendizaje y la resolución de problemas.

Machine Learning (Aprendizaje automático)

Un gran avance en la consecución de la IA

Los algoritmos de aprendizaje automático detectan patrones en grandes conjuntos de datos y aprenden a hacer predicciones procesando datos, en lugar de recibir instrucciones de programación explícitas.

Deep Learning (Aprendizaje profundo)

Una rama avanzada del aprendizaje automático

El aprendizaje profundo utiliza redes neuronales, inspiradas en las formas en que las neuronas interactúan en el cerebro humano, para ingerir datos y procesarlos a través de múltiples iteraciones que aprenden características cada vez más complejas de los datos y hacen predicciones cada vez más sofisticadas.

IA Generativa

Una rama avanzada del aprendizaje profundo

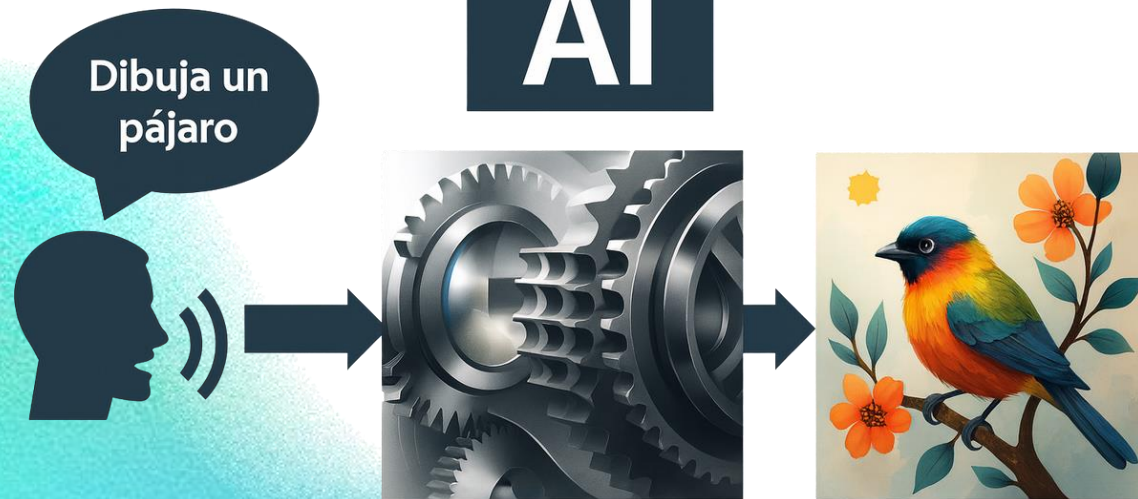
La IA generativa es una rama del aprendizaje profundo que utiliza redes neuronales excepcionalmente grandes llamadas LLM o modelos de lenguaje grandes (con cientos de miles de millones de neuronas) que pueden aprender patrones especialmente abstractos. Los modelos de lenguaje aplicados para interpretar y crear texto, videos, imágenes y datos se conocen como IA generativa.

¿Qué es la

IA?

AI

Un sistema que puede generar contenido nuevo de tipo texto, imagen, sonido o vídeo a partir de una entrada dada.

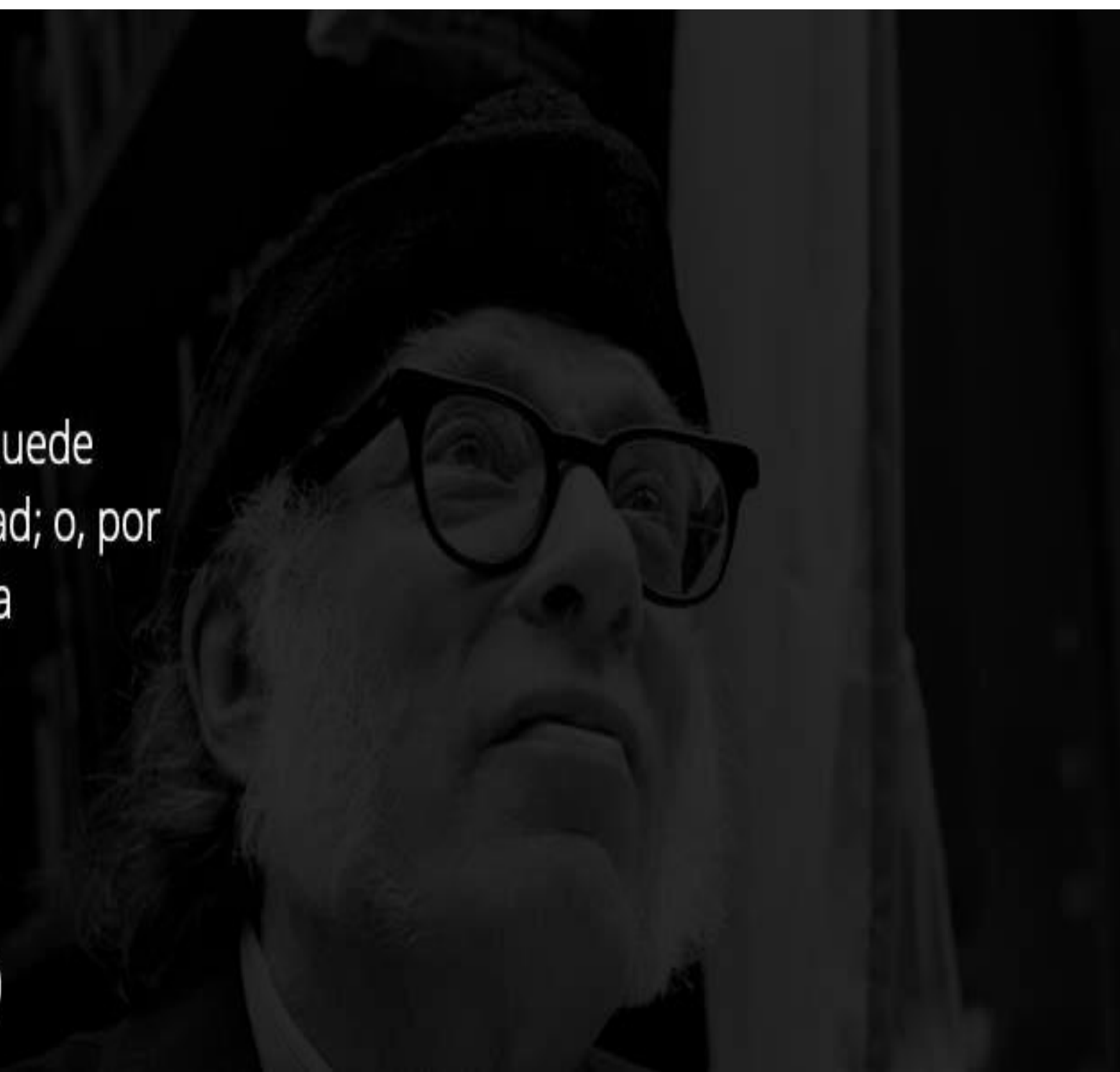


¿Es la pregunta correcta para obtener la respuesta que vemos?

Yo, Robot

"Ninguna máquina puede
dañar a la Humanidad; o, por
inacción, dejar que la
Humanidad sufra
daño"

Asimov, Isaac. (1950)



“Ninguna máquina puede dañar a la Humanidad; o, por inacción, dejar que la Humanidad sufra daño”
Esta es una versión de la **Ley Cero de la Robótica**, propuesta por Isaac Asimov.

¿Qué significa?

- 👉 Una inteligencia artificial **no solo debe evitar hacer daño directamente,**
- 👉 sino que también debe **actuar para evitar daños, incluso si no hacer nada también perjudica.**

Aunque fue escrita en 1950, hoy es extremadamente relevante:

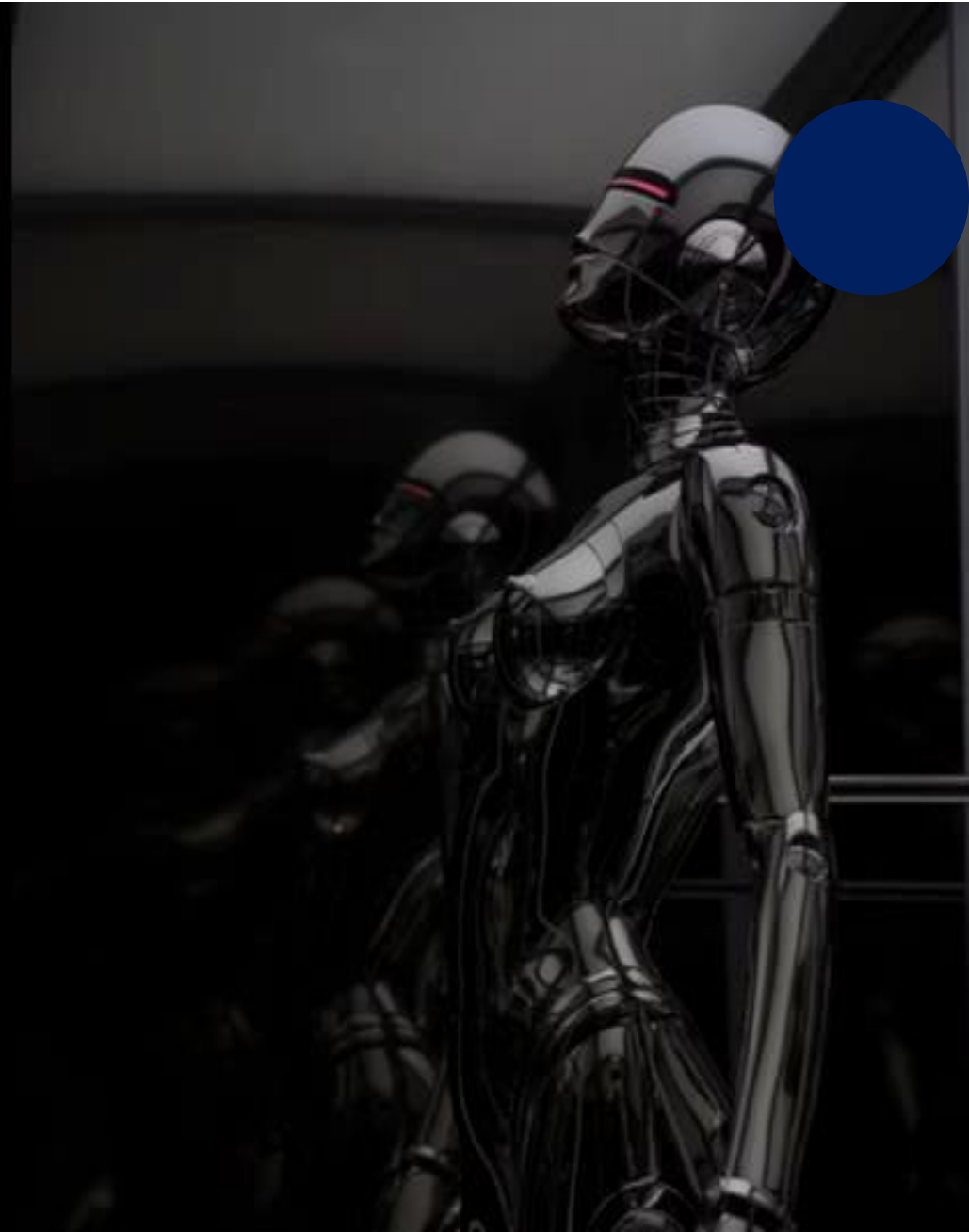
- **1. Decisiones automatizadas**
 - IA en salud, seguridad o transporte
 - ¿Debe actuar o no ante un riesgo?
 - 👉 Ejemplo: un sistema de tráfico inteligente
- **2. Responsabilidad algorítmica**
 - ¿Quién responde si una IA no actúa y ocurre un daño?
 - 👉 Problema clave en regulación actual
- **3. Ética en IA**
 - No basta con “no hacer daño”
 - También implica **prevenir daño**
 - 👉 Esto introduce dilemas:
 - ¿Qué es “el bien de la humanidad”?
 - ¿Quién lo define?

“Asimov plantea que una IA no solo debe evitar hacer cosas malas, sino también evitar que ocurran cosas malas, incluso si no intervenir sería más fácil.”

Definición de Inteligencia

Es una facultad propia de ciertas clases de seres vivos, que les permite pensar, obrar a voluntad, ser conscientes de su existencia, establecer relaciones con el mundo exterior, aprender, ser originales, adaptarse, razonar, autocorregirse, y atender sus necesidades.

By Paul y Marie Christine Haton



Paul Haton y Marie-Christine Haton fueron dos investigadores franceses muy influyentes en los inicios de la **Inteligencia Artificial en Europa**.

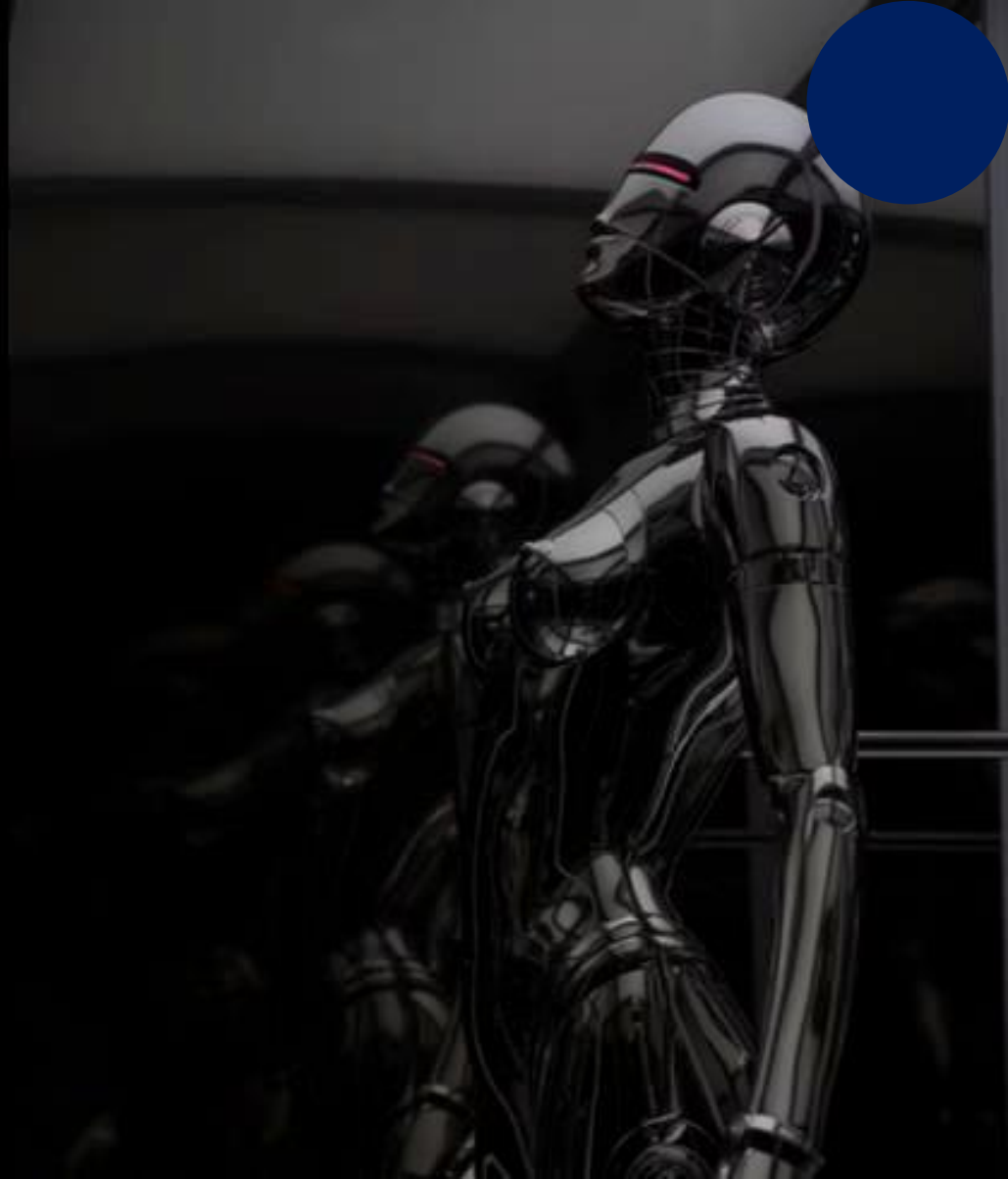
Definición de IA

La capacidad de un artefacto (creado por humanos) de realizar los mismos tipos de funciones que caracterizan a la inteligencia humana, nombrados anteriormente.

La inteligencia artificial es un campo muy amplio, pero su desarrollo se enfoca principalmente en hacer que las máquinas y robots se parezcan cada vez más a los humanos, tanto en la forma de pensar, como en la forma de actuar.

By Paul y Marie Christine Haton

Los Haton representan una etapa de la IA donde el conocimiento se modelaba mediante reglas y lógica, antes del auge del machine learning actual.



Modelo de lenguaje preentrenado de Open AI

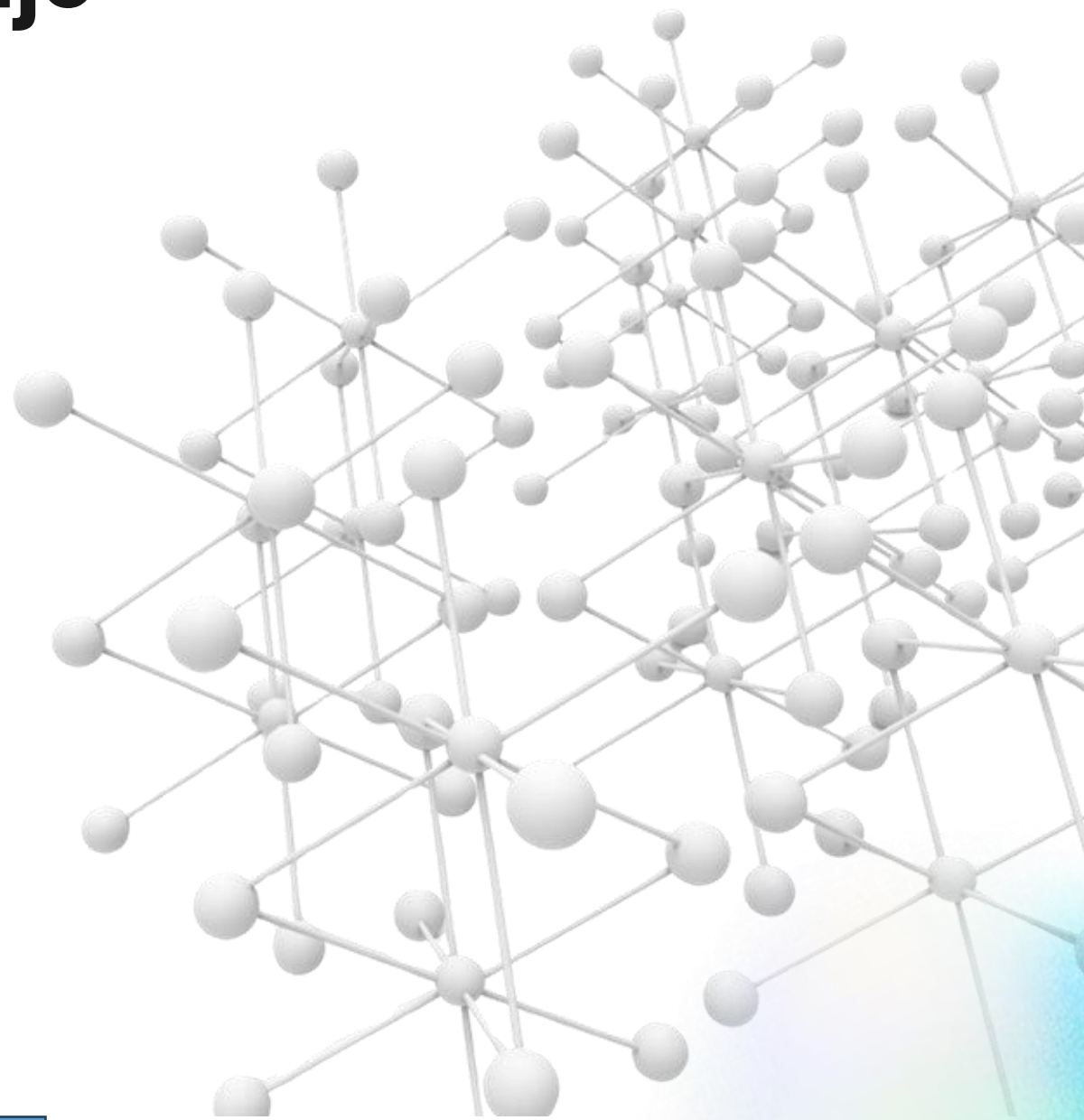
GPT (Generative Pre-Trained
Transformer)

Son modelos de predicción
lingüística basados en redes
neuronales y contruidos sobre
la arquitectura Transformer.

Analizan las consultas en
lenguaje natural, conocidas
como **PROMPTS**, y predicen la
mejor respuesta posible
basándose en su comprensión
del lenguaje.



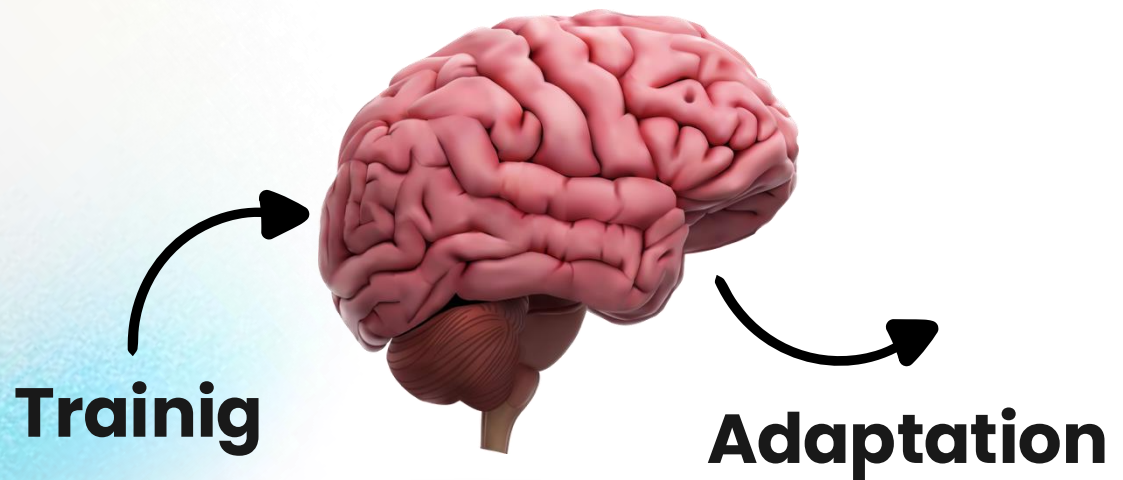
ChatGPT



Large Language Model (LLM)

Un modelo lingüístico grande.

(LLM) es un tipo de programa de inteligencia artificial (IA) que puede reconocer y generar texto, entre otras tareas. Los LLM se capacitan con enormes conjuntos de datos — de ahí el adjetivo "grande".



Acerca de normativas necesarias



Caso: Colombia

Proyecto de Ley de Inteligencia Artificial en Colombia. (43/2025)

- Motivos
- Justificación
- Contexto global
- Importancia de la IA en Colombia
- Necesidad de Regulación
- Normativa Constitucional y Legal (internacionales)
- Objetivo de la Ley
- Impacto social, económico, fiscal, ambiental



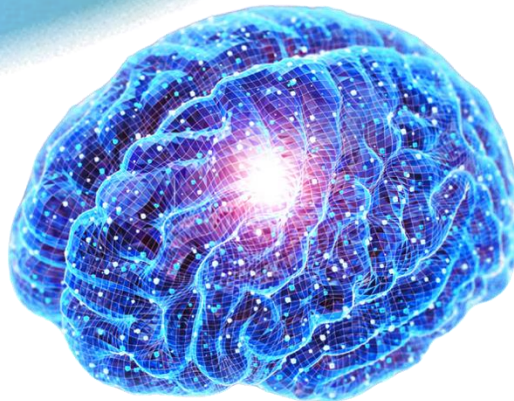
Objetivo del PL (43/2025).

Objetivo general:

Garantizar el desarrollo ético y responsable de la IA en Colombia.

Objetivos específicos:

- 1. Proteger los derechos fundamentales y la seguridad de las personas.
- 2. Fomentar la innovación y la competitividad.
- 3. Asegurar la transparencia y la explicabilidad de los sistemas de IA
- 4. Promover la colaboración interdisciplinaria, interinstitucional e internacional



Avances en leyes ya existen referidas a el uso de la IA en países de mundo. (7 destacados)

- Unión Europea
- EEUU
- China
- Corea del Sur
- Reino Unido
- Argentina
- Colombia
-



Unión Europea

El caso más avanzado

La Ley de Inteligencia Artificial de la Unión Europea (AI Act) es actualmente la regulación más completa y estricta del mundo. Entró en vigor desde 2024 y contempla obligaciones aun mas fuertes que se aplicaran entre el 2026 y 2027.

¿Qué hace? clasifica la IA según niveles de riesgo:

- ✓ Riesgo mínimo → permitido
- ⚠ Riesgo alto → regulado estrictamente
- ✗ Riesgo inaceptable → prohibido

Ejemplos prohibidos:

- “Puntuación social” estilo vigilancia masiva
- Manipulación psicológica peligrosa
- Algunos usos de reconocimiento facial



Exige: transparencia, etiquetado de contenido generado por IA , supervisión humana, evaluación de riesgos

Estados Unidos

Estados Unidos todavía **no** tiene una **ley federal integral única** sobre IA, pero sí:

- Órdenes ejecutivas presidenciales
- Regulaciones sectoriales
- Debates legislativos estatales y federales

Se enfocan especialmente en:

- Seguridad nacional
- Deepfakes (contenido falso)
- Privacidad
- Derechos de autor
- Uso de IA en procesos electorales

Algunos estados como California avanzan más rápido que otros.



China

China tiene uno de los modelos regulatorios más fuertes y centralizados.

China ya regula:

- Algoritmos de recomendación
- IA generativa
- Deepfakes (contenido falso)
- Humanos digitales
- Etiquetado obligatorio de contenido sintético

Además, recientemente tribunales chinos comenzaron a limitar despidos laborales basados únicamente en reemplazo por IA.

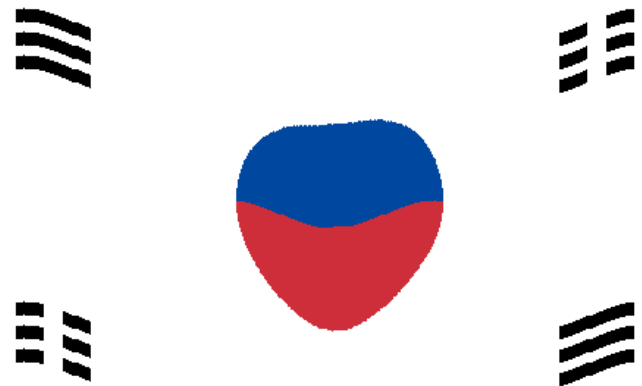


Corea del Sur

Corea del Sur avanzó fuertemente en regulación de deepfakes y contenidos manipulados con IA.

Su enfoque intenta:

- Proteger derechos digitales
- Evitar manipulación
- Permitir innovación tecnológica



Reino Unido

Reino Unido eligió una estrategia más flexible:

- Regulación por sectores
- Menos restricciones iniciales
- Fuerte impulso a innovación

Prefieren adaptar leyes existentes antes que crear una gran ley única.



Argentina

Argentina todavía no posee una ley integral específica de IA como la UE, pero ya existen:

- proyectos legislativos
- debates regulatorios
- análisis sobre privacidad, responsabilidad y uso ético

Mucho se apoya actualmente en:

- protección de datos personales
- derecho del consumidor
- propiedad intelectual
- responsabilidad civil

El tema está creciendo rápidamente en universidades, gobierno y sector privado.



Argentina

“Guía para el uso responsable y ético de herramientas de Inteligencia Artificial Generativa” elaborada para el Consejo Profesional de Abogados y Procuradores de la Provincia de La Rioja.



CONSEJO PROFESIONAL
DE ABOGADOS Y PROCURADORES
DE LA PROVINCIA DE LA RIOJA

GUÍA PARA EL USO RESPONSABLE Y ÉTICO DE HERRAMIENTAS DE IA GENERATIVA

CONSEJO PROFESIONAL DE ABOGADOS Y PROCURADORES
DE LA PROVINCIA DE LA RIOJA



La Inteligencia Artificial Generativa (IAGen) puede mejorar la productividad, precisión y accesibilidad de los servicios jurídicos. Pero su uso sin control humano puede generar consecuencias éticas y disciplinarias para los profesionales del Derecho.



El abogado sigue siendo plenamente responsable de todo contenido generado o asistido por IA.

ANTECEDENTES RELEVANTES (2025)

En Argentina, se registraron sanciones a abogados por usar información falsa o jurisprudencia inexistente generada por IA sin verificación humana.



Rosario (Santa Fe) – 2025

Abogado sancionado por presentar citas jurisprudenciales inexistentes generadas por IA.



Río Negro – 2025

Abogado sancionado por acompañar doctrina falsa creada por IA en un escrito judicial.



Estos casos evidencian los riesgos del uso irresponsable de estas herramientas.

PRINCIPIOS RECTORES PARA EL USO DE IA



1 RESPONSABILIDAD PROFESIONAL

El abogado sigue siendo plenamente responsable de todo contenido generado o asistido por IA.



2 CONTROL HUMANO PERMANENTE

La IA no debe tomar decisiones autónomas sin supervisión profesional.



3 VERACIDAD Y TRANSPARENCIA

Toda información debe verificarse con fuentes oficiales para evitar “alucinaciones” de la IA.



4 CONFIDENCIALIDAD Y PROTECCIÓN DE DATOS

No deben cargarse datos sensibles ni expedientes en plataformas inseguras.



5 CAPACITACIÓN CONTINUA

Los profesionales deben formarse en IA, ciberseguridad y técnicas de prompting.



6 USO PROPORCIONAL Y CONTEXTUAL

La IA debe utilizarse de manera adecuada y justificada según cada situación.

RECOMENDACIONES PRÁCTICAS



Validar toda cita o contenido generado por IA.



Supervisar permanentemente borradores y documentos.



Informar al cliente cuando se utilice IA en trabajos jurídicos.



Proteger el secreto profesional y anonimizar información sensible.



Utilizar únicamente plataformas seguras y confiables.



Promover la formación permanente desde el Consejo Profesional.

ÉTICA Y RÉGIMEN DISCIPLINARIO: PUEDEN CONSIDERARSE FALTAS PROFESIONALES



Presentar documentos con citas falsas o inventadas.



Vulnerar el secreto profesional.



Elaborar escritos judiciales sin revisión humana.



Difundir información errónea o manipulada.

BUENAS PRÁCTICAS CLAVE

- ✓ Usar la IA como asistente, no como reemplazo del criterio profesional.
- ✓ Verificar siempre. No confiar ciegamente en los resultados.
- ✓ Documentar el proceso de trabajo y las fuentes consultadas.
- ✓ Respetar las normas legales, éticas y deontológicas vigentes.
- ✓ Priorizar siempre el interés del cliente y la justicia.



CONCLUSIÓN

La IA es una herramienta transformadora para el ámbito jurídico, pero su implementación debe realizarse bajo criterios de ética, prudencia, responsabilidad profesional, control humano y respeto por las normas legales vigentes.

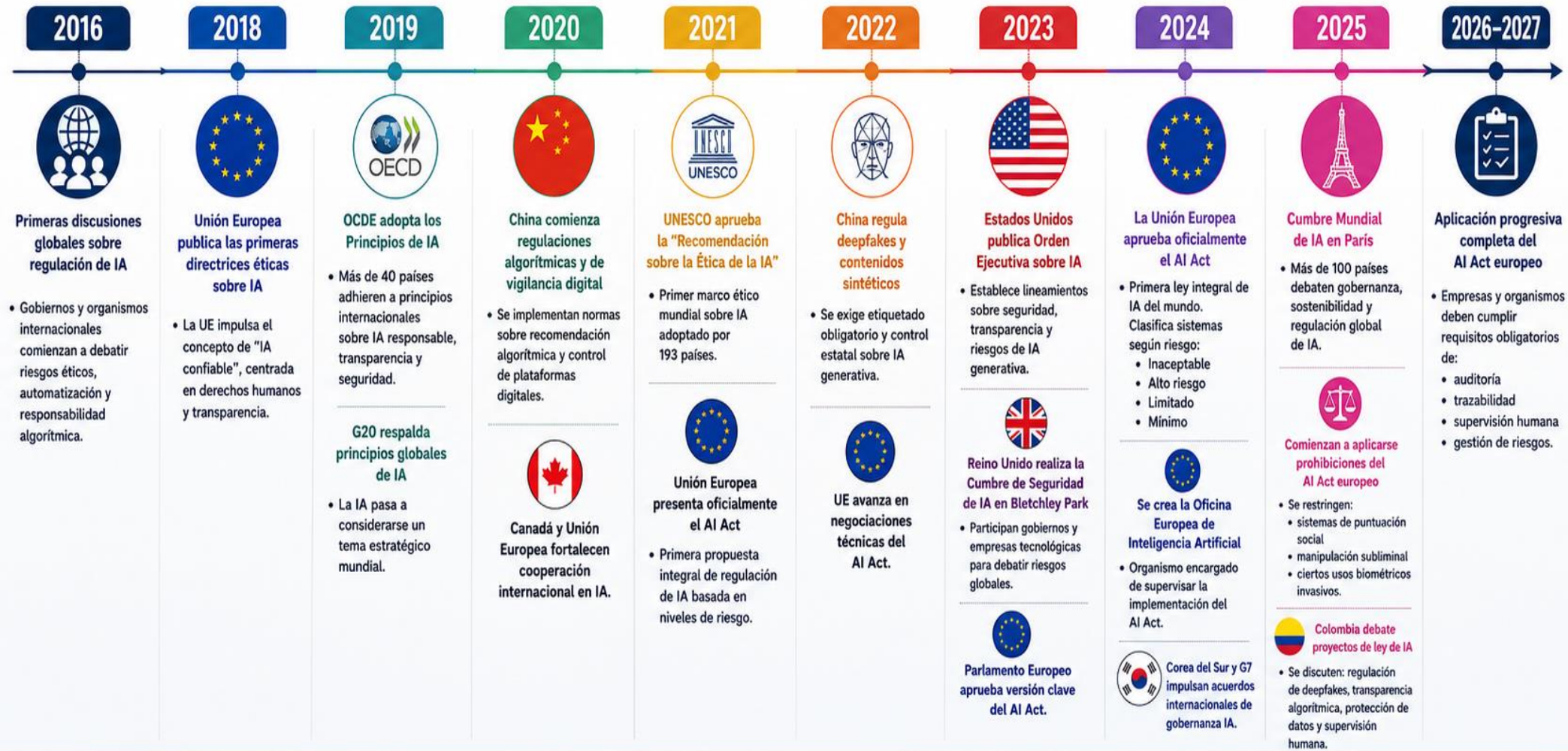


El objetivo es fortalecer la confianza pública en el ejercicio de la abogacía y garantizar un servicio jurídico de calidad, seguro y ético.



LÍNEA DE TIEMPO MUNDIAL DE LEYES Y REGULACIONES DE IA

Hitos clave en la creación de marcos legales y regulaciones para una Inteligencia Artificial ética, segura y centrada en las personas.



TENDENCIA MUNDIAL ACTUAL

La mayoría de los países avanzan hacia regulaciones centradas en:



Ética



Transparencia



Supervisión humana



Protección de datos



Regulación de deepfakes



Seguridad algorítmica



Derechos humanos



IA responsable

REFLEXIÓN FINAL

La historia de las leyes de IA muestra que el desafío ya no es solamente crear inteligencia artificial, sino aprender a gobernarla responsablemente para proteger a las personas, la democracia y el futuro digital de la sociedad.





TENDENCIAS GLOBALES QUE SE REPITEN

Muchos países están empezando a exigir:



1

TRANSPARENCIA

Que las personas sepan cuándo interactúan con IA.



2

ETIQUETADO

Marcas de agua o avisos en imágenes, videos o textos generados artificialmente.



3

SUPERVISIÓN HUMANA

Que decisiones críticas no queden completamente automatizadas.



4

PROTECCIÓN LABORAL

Debates sobre reemplazo laboral, vigilancia de empleados y automatización.



5

PROTECCIÓN FRENTE A DEEPFAKES

Especialmente en:

-  política
-  menores
-  contenido sexual no consentido



La regulación de la IA avanza en todo el mundo con un objetivo común:
INNOVAR CON RESPONSABILIDAD Y PROTEGER A LAS PERSONAS.



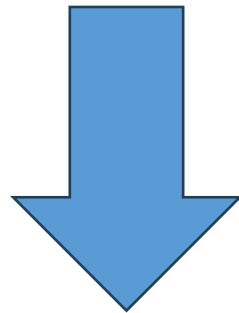
Lo más importante

La regulación de la IA recién está comenzando. Lo que hoy ocurre con la IA se parece mucho a:

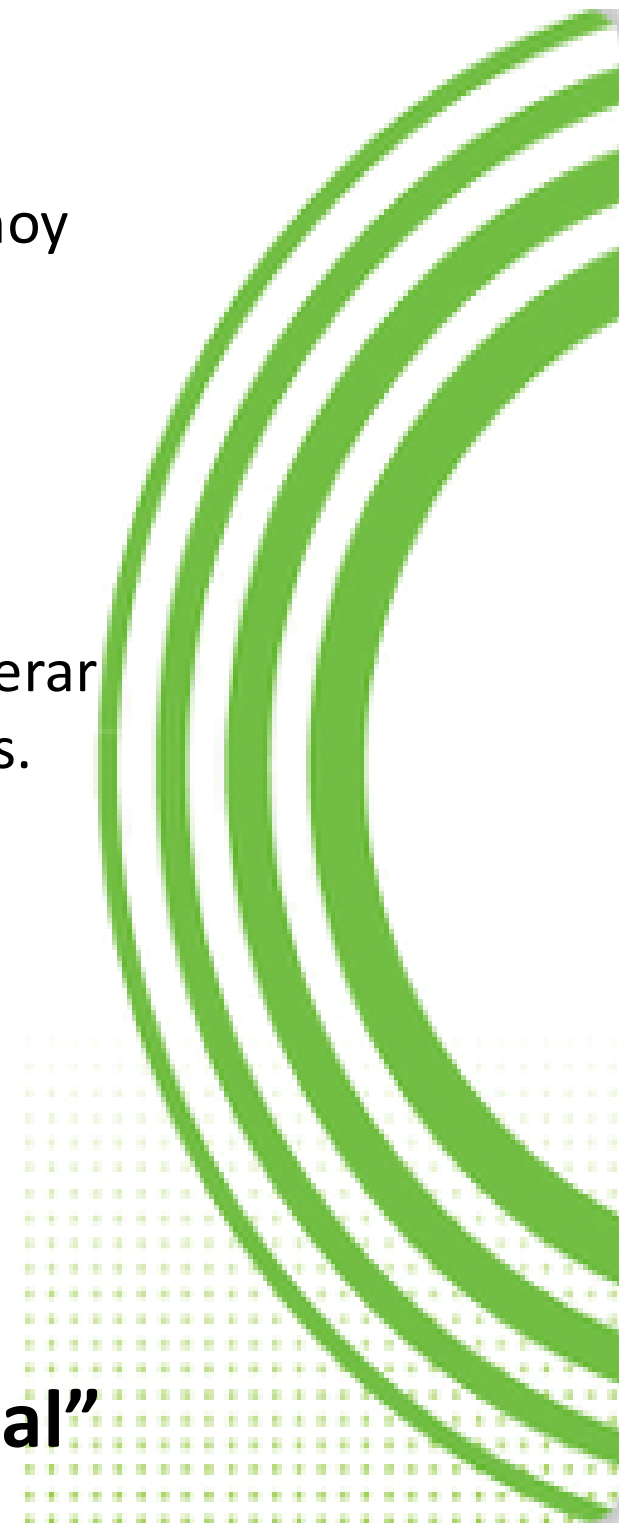
- Los primeros años de Internet
- El inicio de las redes sociales

La diferencia es que los gobiernos entendieron que esperar demasiado puede generar impactos sociales enormes.

Por eso hoy el mundo está entrando en una nueva etapa.....:



“La gobernanza de la Inteligencia Artificial”



La **gobernanza de la Inteligencia Artificial (IA)** es el conjunto de *normas, principios, políticas, instituciones y mecanismos* que buscan asegurar que la IA se desarrolle y utilice de manera:

- Ética
- Segura
- Transparente
- Responsable
- Alineada con los derechos humanos y los valores sociales



Su **objetivo principal** es que la IA *beneficie a las personas y a la sociedad, minimizando riesgos y evitando abusos.*

En términos simples

- La gobernanza de la IA busca responder preguntas como:
- ¿Quién controla cómo se usa la IA?
- ¿Qué límites debe tener?
- ¿Quién es responsable si la IA causa daño?
- ¿Cómo se protege la privacidad?
- ¿Cómo evitar discriminación o manipulación?
- ¿Qué decisiones deben seguir teniendo supervisión humana?

¿Qué incluye?

- La gobernanza de la IA abarca:
- Leyes y regulaciones
- Normas éticas
- Supervisión humana
- Transparencia en algoritmos
- Protección de datos
- Auditorías y controles
- Gestión de riesgos
- Rendición de cuentas
- Seguridad digital
- Políticas públicas

Ejemplos prácticos

- Exigir que un chatbot avise que es IA.
- Regular deepfakes.
- Limitar reconocimiento facial masivo.
- Proteger empleos frente a automatización extrema.
- Exigir explicaciones sobre decisiones algorítmicas.
- Supervisar IA utilizada en salud, justicia o educación.

Principios más importantes

- Transparencia
- Equidad
- Privacidad
- Responsabilidad
- Seguridad
- Inclusión
- Supervisión humana
- Sostenibilidad

¿Por qué es importante?

Porque la IA ya impacta en:

- Educación,
- Empleo,
- Salud,
- Finanzas,
- Política,
- Seguridad,
- Justicia,
- Redes sociales.

**Sin gobernanza adecuada,
pueden aparecer riesgos**

- Discriminación algorítmica,
- Manipulación,
- Vigilancia excesiva,
- Pérdida de privacidad,
- Desinformación,
- Deepfakes,
- Concentración de poder tecnológico.

La gobernanza de la IA no busca frenar la innovación. Busca que la inteligencia artificial evolucione de forma responsable, poniendo a las personas y a la sociedad en el centro del desarrollo tecnológico.

LA GOBERNANZA DE LA INTELIGENCIA ARTIFICIAL



¿QUÉ ES?

La gobernanza de la Inteligencia Artificial es el conjunto de **principios, normas, instituciones y procesos** que orientan, supervisan y regulan el desarrollo, despliegue y uso de sistemas de IA para generar confianza, minimizar riesgos y maximizar beneficios para la sociedad.

Marcos, políticas y prácticas que aseguran que la IA se desarrolle y use de forma **ética, segura, transparente y alineada** con los valores y derechos humanos.



¿POR QUÉ ES IMPORTANTE?



Genera confianza en la IA



Protege los derechos y la dignidad humana



Reduce riesgos y previene daños



Promueve la equidad y la inclusión



Impulsa la innovación responsable



Fortalece la cooperación internacional



PRINCIPIOS CLAVE

- **Ética:** Respeto a los valores humanos.
- **Transparencia:** Información clara y accesible.
- **Equidad:** Evitar sesgos y discriminación.
- **Responsabilidad:** Definir quién responde.
- **Privacidad y protección de datos.**
- **Seguridad:** Robustez y control de riesgos.
- **Rendición de cuentas:** Supervisión y auditoría.
- **Sostenibilidad:** Impacto social y ambiental.

PILARES DE LA GOBERNANZA DE LA IA

1 MARCO NORMATIVO



- Leyes y regulaciones claras.
- Basadas en derechos humanos.
- Adaptables a la evolución tecnológica.

2 INSTITUCIONES Y POLÍTICAS



- Entidades responsables y coordinadas.
- Políticas públicas basadas en evidencia.
- Participación multiactor.

3 GESTIÓN DE RIESGOS



- Evaluación de impactos y riesgos.
- Enfoque basado en el ciclo de vida de la IA.
- Medidas de mitigación y monitoreo continuo.

4 TRANSPARENCIA Y EXPLICABILIDAD



- Sistemas comprensibles y auditables.
- Información clara sobre cómo y por qué la IA toma decisiones.

5 SUPERVISIÓN HUMANA Y CONTROL



- La IA debe apoyar, no reemplazar el juicio humano.
- Decisiones críticas con supervisión humana.

6 INCLUSIÓN Y PARTICIPACIÓN



- Involucrar a diversos grupos sociales.
- Escuchar a las comunidades afectadas.
- Cerrar brechas digitales.

ACTORES INVOLUCRADOS



Gobiernos



Empresas



Academia



Sociedad civil



Ciudadanía

La colaboración entre todos es esencial para una IA justa y beneficiosa.

¿EN QUÉ APLICA?



Salud



Educación



Finanzas



Industria



Justicia



Transporte

...

En todos los sectores donde la IA impacta la vida de las personas.

DESAFÍOS ACTUALES

- Avance tecnológico más rápido que la regulación.
- Falta de estándares comunes globales.
- Sesgos en datos y algoritmos.
- Concentración de poder y desigualdad.
- Desinformación y mal uso (deepfakes, manipulación).
- Brecha de capacidades y talento.



OBJETIVO FINAL

Asegurar que la Inteligencia Artificial esté al servicio del bienestar humano, el desarrollo sostenible y la construcción de sociedades más justas y conscientes.



**UNA IA BIEN GOBERNADA,
ES UNA IA QUE GENERA FUTURO.**

Las **universidades** de todo el mundo ya comenzaron a implementar regulaciones concretas sobre Inteligencia Artificial, especialmente desde la explosión de la IA generativa como [ChatGPT](#), Gemini o Claude.



Lo interesante es que las universidades pasaron rápidamente de:
“prohibir la IA”

a

“regular su uso responsable”.

Hoy el enfoque dominante es:

“La IA puede usarse, pero bajo reglas claras”.

LA REGULACIÓN DE LA IA EN LAS UNIVERSIDADES DEL MUNDO

LA IA PUEDE USARSE, PERO BAJO REGLAS CLARAS



MASSACHUSETTS
INSTITUTE OF
TECHNOLOGY

ENFOQUE

USO PERMITIDO CON TRANSPARENCIA

- ✓ Los estudiantes deben declarar cuándo y cómo usan IA en sus trabajos.



HARVARD
UNIVERSITY

ENFOQUE

USO PERMITIDO CON SUPERVISIÓN HUMANA

- ✓ La IA es una herramienta de apoyo, no un reemplazo del juicio humano.



UNIVERSITY OF
CAMBRIDGE

ENFOQUE

USO PERMITIDO CON RESPONSABILIDAD ÉTICA

- ✓ Políticas claras sobre cuándo, cómo y por qué usar IA.



東京大学
THE UNIVERSITY
OF TOKYO

ENFOQUE

USO PERMITIDO CON NORMAS ESTRICTAS

- ✓ Directrices específicas para cada asignatura y nivel académico.



Stanford
University

ENFOQUE

USO PERMITIDO CON RESPONSABILIDAD

- ✓ Se exige citar herramientas de IA y mantener la integridad académica.



UNIVERSITY OF
OXFORD

ENFOQUE

USO PERMITIDO CON CRITERIO

- ✓ La IA puede apoyar el aprendizaje, pero no reemplaza el pensamiento crítico ni la autoría.



NUS
National University
of Singapore

ENFOQUE

USO PERMITIDO CON GOBERNANZA

- ✓ Marco institucional que regula la IA con enfoque en equidad, privacidad y seguridad.



THE UNIVERSITY OF
MELBOURNE

ENFOQUE

USO PERMITIDO CON TRANSPARENCIA

- ✓ Se debe informar el uso de IA y garantizar trabajos originales.



- 🔍 TRANSPARENCIA
- 🎓 INTEGRIDAD ACADÉMICA
- 🔒 PRIVACIDAD DE DATOS
- 👤 SUPERVISIÓN HUMANA
- 📋 RENDICIÓN DE CUENTAS
- 📊 EVALUACIÓN JUSTA



TENDENCIA
GLOBAL



POLÍTICAS CLARAS
Y PÚBLICAS



FORMACIÓN EN
USO ÉTICO DE IA



PROTECCIÓN DE DATOS
Y PRIVACIDAD



EVALUACIÓN JUSTA
Y EQUITATIVA



LA IA PUEDE USARSE,
PERO BAJO REGLAS
CLARAS

Principales regulaciones universitarias que ya existen:

- Políticas de integridad académica
- Obligación de declarar uso de IA
- Sistemas “Semáforo”
- Reformulación de evaluaciones
- Limitaciones a detectores de IA (NO son confiables)
- Protección de privacidad y datos
- Capacitación obligatoria en IA
- Regulaciones para investigación científica



MIT Professional Education

Política de IA

Uso de herramientas de IA generativa: Este curso reconoce el potencial de las herramientas de inteligencia artificial (IA) generativa para el aprendizaje y la creatividad de los participantes. Se anima a los participantes a utilizar herramientas de IA para apoyar sus actividades, cuando sea apropiado, asegurándose de hacerlo de manera ética y responsable. Algunas actividades incluso pueden integrar herramientas de IA. Es importante que:

Usar herramientas de IA solo cuando sea apropiado para su nivel de aprendizaje y comprensión. Las herramientas de IA no deben utilizarse para sustituir su propio razonamiento o análisis, ni para evitar el compromiso con el contenido del curso. Explicar y proporcionar evidencia de cómo utilizó la herramienta de IA y cómo contribuyó a su actividad. Explique qué aprendió de la herramienta, cómo verificó su precisión y confiabilidad, cómo integró el resultado en su trabajo y cómo reconoció sus limitaciones y sesgos.

Al trabajar en grupo, discuta el uso de herramientas de IA con los demás integrantes y lleguen a un consenso sobre cómo planean utilizarla, explicando y evidenciando cómo la herramienta contribuyó a la actividad.

Citar adecuadamente cualquier herramienta de IA utilizada. Por ejemplo, las referencias en formato APA exigen incluir el nombre de la herramienta de IA, la fecha de acceso, la URL de la interfaz y el comando o consulta específicos utilizados para generar el resultado. Por ejemplo:

ChatGPT. “¿Cómo puedo implementar una estrategia de sostenibilidad en mi empresa?” Accedido el 8 de mayo de 2024.

<https://chat.openai.com>

Para más información sobre cómo citar y referenciar herramientas de IA generativa de manera consistente, consulte la siguiente página:

<https://libguides.mit.edu/cite-AI-tools>

Asumir la plena responsabilidad de cualquier error o equivocación cometido por la herramienta de IA. Siempre verifique y edite el resultado antes de enviar su trabajo.

Debe usar herramientas de IA de manera ética y responsable. Copiar o parafrasear contenido generado sin citación o evidencia, utilizar los resultados como si fueran su propio trabajo, sin verificación, o emplear los resultados para tergiversar sus conocimientos o habilidades se considera deshonestidad académica. Estas prácticas resultarán en una calificación de cero para la actividad y posibles sanciones disciplinarias. Si tiene dudas sobre qué constituye el uso ético y responsable de herramientas de IA, consulte con la persona facilitadora de aprendizaje antes de enviar su trabajo.

No es cualquier tarea... ¡Es un desafío para poner a prueba sus ideas!



¡ATENCIÓN, EQUIPO! TIENEN UN **RETO**

SUGERIDO POR EL DOCENTE

Usen su creatividad, pensamiento crítico y trabajo en equipo. ¡Ustedes pueden!



¿EN QUÉ CONSISTE?

Resolver un desafío real o planteado, aplicando lo aprendido, investigando, analizando y proponiendo soluciones originales y viables.

¡ACEPTEN EL DESAFÍO!

1 ENTIENDAN EL RETO

Lean con atención las instrucciones y el problema a resolver.



¿Qué se pide?
¿Cuál es el objetivo?
¿Qué sabemos?

2 INVESTIGUEN Y ANALICEN

Busquen información confiable, hagan preguntas y analicen diferentes puntos de vista.



Datos + Ideas =
Mejores soluciones

3 PROPONGAN SU SOLUCIÓN

Diseñen, creen y desarrollen su propuesta. ¡Sean creativos y prácticos!



Piensen diferente,
¡marquen la diferencia!

4 PRESENTEN SU TRABAJO

Organicen su información y presenten su solución de forma clara, visual y convincente.



Comunicar bien
también es parte
del reto.

5 REFLEXIONEN Y MEJOREN

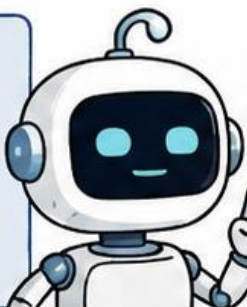
Evalúen lo aprendido, los resultados y cómo pueden mejorar.



Cada reto es una
oportunidad para
crecer.

RECUERDEN

- ✓ Trabajo en equipo
- ✓ Respeto y escucha
- ✓ Responsabilidad
- ✓ Originalidad
- ✓ Ética académica



¡No se trata de tener todas las respuestas, sino de tener ganas de encontrarlas!



EL VERDADERO PREMIO

- 🧠 Aprender algo nuevo
- 🚀 Desarrollar habilidades
- 👥 Trabajar juntos
- ★ Sentirse orgullosos de lo logrado



¡Acepten el reto y conviértanse en protagonistas de su aprendizaje!

Retos: segundo módulo

Elige una IA y escribe tres PROMPTS, desarrollando brevemente como y porque lo escribiste.

Puedes identificar como el Objetivo del Proyecto de Ley número 43/2025 de Colombia, influirá en las acciones a futuro de la IA, en tu contexto social en general y la gobernanza de la IA . Ejemplifica y desarrolla tu respuesta.

¿Cual crees será la tendencia dominante en las instituciones educativas con la mirada colocada en el 2030?

GRACIAS