
CONCEPTOS BÁSICOS DE INTELIGENCIA ARTIFICIAL PARA DEEPSEEK

Para entender cómo funciona DeepSeek, primero es importante conocer los principios básicos de la inteligencia artificial. En esta sección explicaremos qué es la IA, cómo se diferencia del aprendizaje automático y las redes neuronales, y cómo DeepSeek interpreta nuestras preguntas mediante el procesamiento del lenguaje natural. Además, veremos los módulos clave que hacen posible su funcionamiento y reflexionaremos sobre la ética y la transparencia en su desarrollo.

2.1 ¿QUÉ ES LA INTELIGENCIA ARTIFICIAL?

La Inteligencia Artificial (IA) es un área de la tecnología que busca crear programas o máquinas capaces de realizar tareas que, hasta hace poco, solo se creía que podían hacer los humanos. Estas tareas incluyen reconocer voces, entender textos o tomar decisiones basadas en análisis de datos. Hoy en día, la IA se encuentra en aplicaciones tan comunes como los asistentes virtuales de los teléfonos, los sistemas de recomendación de películas y los filtros de correo no deseado.



2.1.1 Definición simple y usos generales

Si lo explicamos de forma sencilla, la IA es como darle “inteligencia” a un software para que observe el entorno, aprenda de la experiencia y actúe con base en lo que ha aprendido. Gracias a ello, puede reconocer patrones, procesar información y encontrar soluciones a problemas de manera más rápida y eficiente que un programa tradicional. Por ejemplo, en la medicina se emplea para analizar radiografías y detectar signos de enfermedades, en los coches para la conducción autónoma y en las redes sociales para moderar contenido inapropiado. De esta manera, la IA aporta beneficios en distintos ámbitos como la salud, la industria, el entretenimiento y la educación, entre otros. A continuación, se presentan algunos de sus usos más importantes en la actualidad:

Asistentes virtuales y chatbots

Los asistentes de voz como Siri, Alexa o Google Assistant utilizan IA para responder preguntas, programar recordatorios y controlar dispositivos inteligentes. Además, los chatbots en páginas web y redes sociales ayudan a resolver dudas de los clientes, ofreciendo respuestas rápidas sin necesidad de intervención humana.

Reconocimiento de voz y procesamiento del lenguaje natural

La IA permite que los dispositivos comprendan y transcriban lo que decimos. Esto se usa en aplicaciones de dictado, subtítulos automáticos y sistemas de traducción, facilitando la comunicación en diferentes idiomas y mejorando la accesibilidad para personas con discapacidad auditiva.

Recomendaciones personalizadas

Plataformas como Netflix, Spotify y YouTube utilizan IA para sugerir contenido basado en los gustos y hábitos de los usuarios. Lo mismo ocurre con tiendas en línea como Amazon, que analizan el historial de compras para ofrecer productos que pueden interesar al cliente.

Seguridad informática y detección de fraudes

Los sistemas de IA analizan grandes volúmenes de datos para detectar patrones sospechosos en transacciones bancarias o intentos de ciberataques. Esto ayuda a prevenir fraudes en tarjetas de crédito y protege la información personal de los usuarios en internet.

Diagnóstico médico y salud

En el sector de la salud, la IA se usa para analizar radiografías, identificar síntomas en imágenes médicas y predecir posibles enfermedades basándose en el historial clínico de los pacientes. También se emplea en la investigación de nuevos medicamentos y en la gestión de hospitales para optimizar recursos.

Vehículos autónomos y navegación

Los coches sin conductor, como los desarrollados por Tesla, utilizan IA para interpretar señales de tráfico, detectar obstáculos y tomar decisiones en tiempo real. Además, aplicaciones como Google Maps o Waze analizan el tráfico en directo y sugieren rutas más rápidas.

Automatización en la industria y la robótica

Las fábricas han integrado robots con IA para realizar tareas repetitivas, como ensamblar productos o controlar la calidad de estos. Esto aumenta la eficiencia y reduce los errores humanos en la producción.

Inteligencia Artificial en la educación

Las plataformas de aprendizaje en línea utilizan IA para personalizar el contenido según el nivel y el ritmo de cada estudiante. También ayudan a corregir exámenes automáticamente y a generar explicaciones detalladas de los errores cometidos.

Predicción del clima y gestión del medio ambiente

Los sistemas de IA analizan datos meteorológicos para predecir fenómenos como tormentas o huracanes con mayor precisión. También ayudan a optimizar el consumo de energía en edificios inteligentes y a mejorar la gestión de recursos naturales.

Creación de contenido y generación de imágenes

Los modelos de IA pueden generar texto, música, arte e incluso vídeos de manera automatizada. Por ejemplo, algunas herramientas permiten escribir artículos, componer melodías o diseñar imágenes realistas a partir de descripciones textuales.

Finanzas y análisis de inversiones

Las instituciones bancarias y las plataformas de inversión utilizan IA para analizar mercados financieros y predecir tendencias. Los algoritmos de trading automatizado pueden tomar decisiones en milisegundos basándose en patrones de datos históricos y noticias económicas. También ayudan a los bancos a evaluar riesgos de crédito al analizar el historial financiero de los clientes.

Control de calidad en la industria

En fábricas y líneas de producción, la IA se usa para inspeccionar productos y detectar defectos en tiempo real. Gracias a cámaras de alta resolución y modelos de aprendizaje automático, las máquinas pueden identificar errores en piezas y corregirlos sin necesidad de intervención humana, reduciendo el desperdicio de materiales.

Agricultura inteligente

La IA ayuda a mejorar la productividad agrícola mediante drones y sensores que analizan el estado del suelo y los cultivos. También permite predecir plagas

o enfermedades en las plantas y optimizar el riego con base en datos climáticos, reduciendo el consumo de agua y mejorando los rendimientos de las cosechas.

Inteligencia artificial en la música y el arte

Los sistemas de IA pueden componer canciones, crear ilustraciones o generar textos de manera autónoma. Herramientas como DALL·E, Stable Diffusion o ChatGPT permiten a los artistas y creadores generar contenido original basándose en ideas o descripciones textuales. También se utilizan en la restauración de obras de arte antiguas.

Control de tráfico y planificación urbana

Las ciudades inteligentes emplean IA para mejorar la movilidad, ajustando semáforos en función del tráfico en tiempo real. También se usan modelos predictivos para planificar infraestructuras de transporte y reducir la congestión vehicular.

Supervisión y análisis de redes sociales

Plataformas como Facebook, Twitter o Instagram usan IA para moderar contenido, detectar discursos de odio o eliminar noticias falsas. También analizan tendencias y sugieren publicaciones en función de los intereses de cada usuario.

Búsqueda de empleo y selección de personal

Las empresas utilizan IA para filtrar currículums y encontrar candidatos adecuados para una vacante. También hay asistentes virtuales que pueden realizar entrevistas preliminares, analizando el tono de voz y las respuestas de los postulantes.

Inteligencia Artificial en el deporte

Los entrenadores y atletas utilizan IA para analizar el rendimiento en tiempo real. Cámaras y sensores registran movimientos y generan datos que ayudan a mejorar técnicas en deportes como el fútbol, el tenis o el atletismo. También se emplea para predecir lesiones y optimizar la recuperación de los deportistas.

Detección de terremotos y desastres naturales

Sistemas avanzados de IA analizan patrones sísmicos y datos geológicos para predecir terremotos o alertar sobre posibles tsunamis. Estos sistemas pueden salvar vidas al emitir advertencias tempranas y permitir evacuaciones a tiempo.

Análisis de ADN y genética

En biotecnología, la IA se usa para analizar el ADN y encontrar relaciones entre genes y enfermedades. Esto permite personalizar tratamientos médicos y desarrollar terapias más eficaces para enfermedades genéticas o cáncer.

Videojuegos y entretenimiento interactivo

Los videojuegos modernos incorporan IA para mejorar la experiencia del jugador. Desde personajes no jugables (NPCs) con comportamientos realistas hasta la generación automática de mundos y misiones, la IA crea experiencias más inmersivas y dinámicas.

Seguridad y videovigilancia inteligente

Los sistemas de seguridad usan IA para detectar movimientos sospechosos en cámaras de vigilancia. Algunos incluso pueden identificar rostros y alertar sobre actividades inusuales en tiempo real, mejorando la protección en lugares públicos y privados.

Energía y sostenibilidad

Las compañías eléctricas utilizan IA para optimizar el consumo de energía y gestionar redes inteligentes. También se emplea en plantas solares y eólicas para mejorar la eficiencia y predecir la producción energética en función de las condiciones climáticas.

Administración pública y servicios gubernamentales

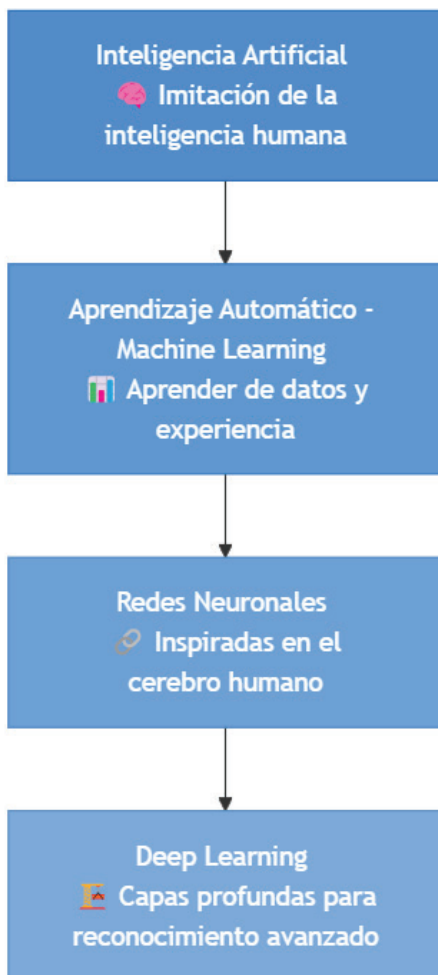
Los gobiernos están implementando IA para mejorar la atención al ciudadano, agilizar trámites burocráticos y detectar fraudes fiscales. También se usa en la planificación de políticas públicas basadas en el análisis de datos socioeconómicos.

Exploración espacial

Las agencias espaciales como la NASA y SpaceX emplean IA para analizar datos del espacio, planificar misiones y controlar vehículos autónomos en Marte. También se usa en telescopios avanzados para procesar imágenes y descubrir nuevos planetas o galaxias.

2.1.2 Diferencias entre IA, aprendizaje automático y redes neuronales

Dentro de la IA existen diferentes enfoques, siendo uno de los más relevantes el **aprendizaje automático (Machine Learning)**, que permite a las máquinas mejorar su rendimiento analizando grandes volúmenes de datos y detectando patrones sin necesidad de seguir instrucciones predefinidas. Un ejemplo común es el filtro de spam en los correos electrónicos, que con el tiempo aprende a distinguir mejor entre mensajes legítimos y no deseados. Dentro del aprendizaje automático, un método especialmente avanzado es el de las **redes neuronales**, que imitan la estructura del cerebro humano a través de capas de nodos interconectados. Las redes neuronales profundas (**Deep Learning**) utilizan múltiples capas para analizar información de manera progresiva, permitiendo grandes avances en reconocimiento de voz, visión por computadora y traducción automática, entre otras aplicaciones.



- **IA (Inteligencia Artificial):** es el concepto general que incluye cualquier sistema que tenga la capacidad de imitar comportamientos “inteligentes” humanos, ya sea aprendiendo, razonando o resolviendo problemas. Dentro de la IA, encontramos varios enfoques y métodos.
 - **Aprendizaje Automático (Machine Learning):** se refiere a la forma en que las máquinas “aprenden” de la experiencia. En lugar de seguir una lista de reglas fijas, estos sistemas analizan grandes cantidades de datos y detectan patrones para mejorar sus respuestas a lo largo del tiempo. Un ejemplo es un filtro de spam que, tras revisar miles de correos, va puliendo su capacidad para separar lo que es correo normal de lo que no lo es.
 - **Redes Neuronales:** son uno de los métodos más populares dentro del aprendizaje automático. Se inspiran en la forma en que funciona el cerebro humano, usando “capas” de nodos (o neuronas artificiales) que procesan la información de forma jerárquica.
 - **Las redes neuronales profundas (Deep Learning) utilizan** muchas capas para reconocer patrones cada vez más complejos, lo que ha permitido avances importantes en áreas como el reconocimiento de voz, la visión por computadora y la traducción automática.

El **Aprendizaje Automático** es una rama de la Inteligencia Artificial que se basa en la idea de que las computadoras pueden identificar patrones en grandes cantidades de datos y, a partir de ellos, “aprender” para realizar tareas sin estar programadas de manera explícita. A diferencia de un software tradicional, donde se indican todas las reglas paso a paso, en Machine Learning el sistema encuentra esas reglas por sí mismo al analizar ejemplos y extraer conclusiones que luego aplica a información nueva.

Para que un algoritmo de Machine Learning funcione, necesita datos de calidad. Estos datos suelen estar organizados en registros o conjuntos de ejemplos que contienen la información relevante para la tarea que se quiere realizar. Por ejemplo, si un modelo tiene que predecir el precio de las casas, necesitará datos sobre el tamaño de la vivienda, el número de habitaciones, la ubicación y otros factores que influyen en su valor. Antes de entrenar el modelo, es habitual limpiar los datos para eliminar valores erróneos, duplicados o inconsistentes, y así asegurarse de que el algoritmo obtenga una visión lo más precisa posible de la realidad.

Una vez que se cuentan con los datos adecuados, el siguiente paso es elegir el tipo de aprendizaje que se va a emplear. Existen varias categorías principales en Machine Learning. Una de ellas es el aprendizaje supervisado, donde se trabaja con datos ya etiquetados (por ejemplo, imágenes de perros y gatos con su correspondiente

etiqueta). Con estos datos, el algoritmo aprende a reconocer qué características diferencian a un perro de un gato, y luego puede clasificar nuevas imágenes que no había visto antes. En el aprendizaje no supervisado, el sistema intenta descubrir patrones por su cuenta, sin etiquetas previas. Esto se utiliza mucho para segmentar clientes o detectar comportamientos anormales sin tener que definir qué es “anormal” desde el principio. Existe también el aprendizaje por refuerzo, en el que el algoritmo aprende mediante la interacción con un entorno, recibiendo recompensas o penalizaciones según sus acciones (muy útil en el desarrollo de robots o sistemas de recomendación).

El proceso de entrenamiento consiste en alimentar al modelo con ejemplos de entrada (los datos) y la salida esperada (si se trata de un aprendizaje supervisado). El algoritmo ajusta internamente sus parámetros para minimizar los errores y mejorar su capacidad de predicción. Estos parámetros pueden ser coeficientes en un modelo de regresión, pesos sinápticos en una red neuronal o reglas de clasificación en un árbol de decisión, dependiendo del tipo de algoritmo. Con cada iteración, se reduce la diferencia entre la predicción y la respuesta real, lo que se conoce como error. Una vez finalizado el entrenamiento, es esencial probar el modelo con datos nuevos o un conjunto de validación para comprobar si ha aprendido de forma adecuada o si está sobreajustado (lo que significa que aprendió demasiado los detalles de los datos de entrenamiento y no generaliza bien).



En la práctica, existen múltiples algoritmos y técnicas de Machine Learning. Algunos son más sencillos, como la regresión lineal o los árboles de decisión, útiles cuando se necesita un modelo rápido y comprensible. Otros, como las redes neuronales profundas, son capaces de detectar patrones extremadamente complejos en voz, texto o imágenes, aunque requieren más recursos de computación y pueden ser menos fáciles de interpretar. La elección del algoritmo depende de varios factores: el tipo de datos, la complejidad del problema y el nivel de precisión que se busca.

Además de los algoritmos y del entrenamiento en sí, el Aprendizaje Automático implica un seguimiento constante. Una vez que un modelo se pone en producción (por ejemplo, en una aplicación que recomienda productos), se necesita revisar su rendimiento y actualizarlo con datos recientes para que no se quede desfasado. Con el paso del tiempo, las condiciones pueden cambiar, y un modelo que funcionaba bien en un momento dado puede ir perdiendo precisión si no se mantiene al día.

Esquema del proceso de Aprendizaje Automático en un contexto de predicción de precios de casas

1. RECOLECCIÓN DE DATOS

- Se recopilan 5.000 registros de viviendas de una ciudad específica.
- Cada registro incluye información como tamaño en metros cuadrados (por ejemplo, entre 50 y 300 m²), número de habitaciones (1 a 5), ubicación (código postal) y antigüedad de la vivienda (en años).
- El precio de venta de cada casa (entre 80.000 y 500.000 euros) se registra como la variable que se desea predecir.

2. LIMPIEZA Y PREPROCESAMIENTO

- Se eliminan los registros con información incompleta o claramente errónea (por ejemplo, números de metros cuadrados extremadamente altos o negativos).
- Se convierten las variables de texto (código postal) en representaciones numéricas o “one-hot encoding” si se requiere que el modelo interprete distintas zonas geográficas.
- Se normalizan valores muy dispersos, como los metros cuadrados, para que el rango de datos quede más equilibrado (por ejemplo, a un rango de 0 a 1).

3. DIVISIÓN DEL CONJUNTO DE DATOS

- Se separan los 5.000 registros en tres subconjuntos:
 - Entrenamiento (70%): datos usados para ajustar los parámetros del modelo.
 - Validación (15%): datos para comprobar la precisión del modelo y evitar que se ajuste en exceso al conjunto de entrenamiento.
 - Prueba (15%): datos finales para comprobar la calidad del modelo después de todos los ajustes.

4. SELECCIÓN DEL MODELO DE APRENDIZAJE AUTOMÁTICO

- Se opta por un modelo de regresión lineal múltiple como primer paso, dado que es sencillo de interpretar.
- También se considera un modelo de bosque aleatorio (Random Forest) para comparar su rendimiento y ver si mejora la precisión.
- Cada modelo se entrena por separado con los mismos datos de entrenamiento para comparar cuál ofrece mejor resultado.

5. ENTRENAMIENTO DE LOS MODELOS

- **Regresión lineal múltiple:** el modelo ajusta coeficientes para cada variable (metros cuadrados, número de habitaciones, etc.).
- **Bosque aleatorio:** se crean múltiples árboles de decisión, cada uno entrenado en una parte de los datos. La predicción final se obtiene promediando las salidas de todos los árboles.
- Durante el entrenamiento, se mide la diferencia (error) entre la predicción del modelo y el precio real de cada casa.

6. EVALUACIÓN INTERMEDIA CON EL CONJUNTO DE VALIDACIÓN

- Se calculan métricas de rendimiento (por ejemplo, Error Cuadrático Medio, RMSE) para ver qué tan cerca están las predicciones de los precios reales.
- Si un modelo presenta un error más bajo, se analiza si puede mejorarse ajustando parámetros (hiperparámetros en el caso del Bosque aleatorio).

7. AJUSTE DE HIPERPARÁMETROS

- Para el modelo de Bosque aleatorio, se prueban diferentes valores de hiperparámetros (como la cantidad de árboles o la profundidad máxima de cada árbol) para encontrar la mejor combinación.
- Se vuelve a entrenar el modelo con estos parámetros óptimos y se verifica nuevamente el rendimiento en el conjunto de validación.

8. EVALUACIÓN FINAL CON EL CONJUNTO DE PRUEBA

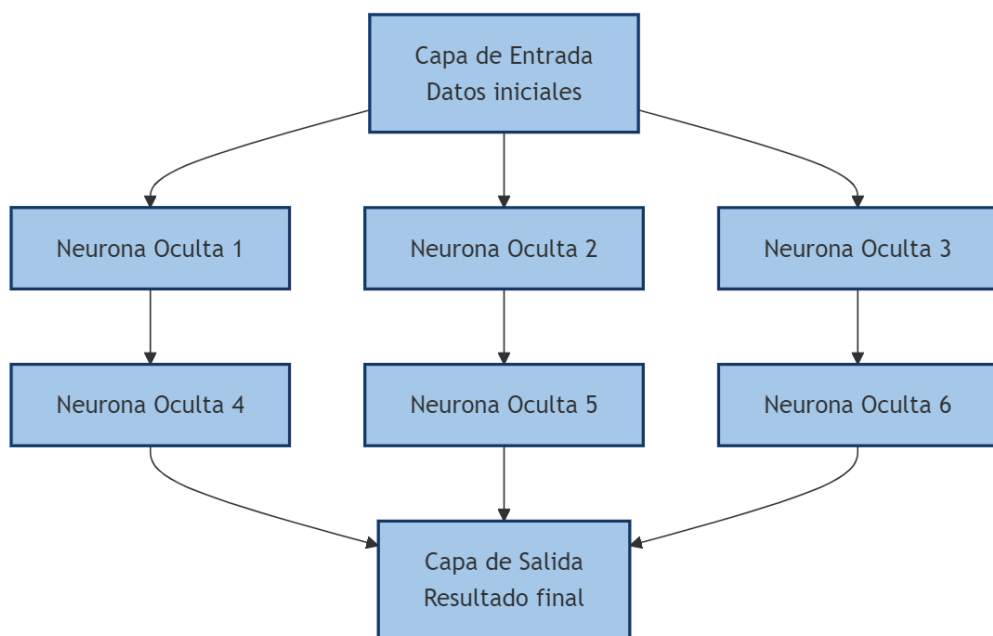
- Una vez seleccionados y ajustados los modelos, se prueban con el 15% de datos que no se usaron ni en entrenamiento ni en validación.
- Se comparan los resultados finales y se elige el modelo que ofrece mejores predicciones en precios de casas.
- Este modelo es el que se implementa, por ejemplo, en una aplicación o en un sistema de recomendaciones para valoraciones inmobiliarias.

9. MANTENIMIENTO Y ACTUALIZACIÓN

- Con el paso del tiempo, los precios de las casas pueden cambiar por factores económicos o nuevas tendencias.
- Periódicamente, se recopilan nuevos registros de viviendas para volver a entrenar o actualizar el modelo y mantenerlo en sintonía con la realidad del mercado.

Las **redes neuronales** funcionan inspirándose en el cerebro humano, utilizando **neuronas artificiales** organizadas en capas que procesan información y aprenden patrones a partir de los datos. Su principal ventaja es que pueden detectar relaciones complejas en grandes volúmenes de información sin que sea necesario programar reglas específicas.

Una **red neuronal artificial** está compuesta por varias capas de nodos (neuronas artificiales):



- **Capa de entrada:** recibe los datos iniciales y los distribuye a la red.
- **Capas ocultas:** procesan la información a través de múltiples transformaciones matemáticas.
- **Capa de salida:** entrega el resultado final del modelo.

Cada conexión entre neuronas tiene un **peso**, que indica la importancia de una entrada en la predicción final. Estos pesos se ajustan durante el **entrenamiento** de la red.

El proceso de funcionamiento es el siguiente:

- **PASO 1. Entrada de datos**

- Los datos de entrada (imágenes, texto, valores numéricos, etc.) se representan en forma de números y se ingresan a la **capa de entrada**.
- Por ejemplo, en una red que reconoce imágenes de gatos y perros, los píxeles de la imagen se convierten en valores numéricos antes de procesarse.

- **PASO 2. Propagación hacia adelante (Forward Propagation)**

- Los datos avanzan por la red, pasando de una neurona a otra. En cada capa oculta, se aplican **operaciones matemáticas**, como:
 - **Multiplicación por pesos** (ajustan la importancia de cada dato).
 - **Aplicación de una función de activación** (determina si una neurona debe activarse).
- Por ejemplo, en una red de traducción automática, una palabra en español se transforma en varias representaciones numéricas antes de obtener su equivalente en inglés.

- **PASO 3. Cálculo del error**

- Al final, la red genera una **salida**, que se compara con la respuesta real. Se mide la diferencia entre la predicción y el resultado correcto con una función llamada **función de pérdida**.
- Por ejemplo, si la red predice que una imagen es un perro con un 70% de certeza, pero en realidad es un gato, se mide cuánto se equivocó la red.

- **PASO 4. Ajuste de pesos (Backpropagation)**

- Para mejorar sus predicciones, la red ajusta sus pesos usando un método llamado **retropropagación del error (backpropagation)**.
 - Se calcula cuánto influyó cada peso en el error.
 - Se ajustan los pesos usando un algoritmo de optimización como el **Descenso del Gradiente**.
 - Se repite el proceso miles de veces hasta que la red aprende correctamente.
- Por ejemplo, en reconocimiento de voz, la red va mejorando con cada iteración hasta distinguir mejor las palabras.

• PASO 5. Predicción con datos nuevos

- Una vez entrenada, la red puede hacer predicciones sobre nuevos datos que no había visto antes.
- Por ejemplo, un asistente de voz como Siri o Google Assistant entiende nuevas frases gracias al aprendizaje previo.

2.1.3 Tipos de Redes Neuronales

- Dependiendo de la tarea, se pueden usar diferentes tipos de redes neuronales:
 - **Perceptrón multicapa (MLP):** para datos tabulares y problemas simples de clasificación y regresión.
 - **Redes Convolucionales (CNN):** para procesamiento de imágenes y reconocimiento facial.
 - **Redes Recurrentes (RNN):** para texto, voz y series temporales.
 - **Transformers:** utilizadas para el procesamiento de lenguaje natural.

Dentro del Aprendizaje Automático, las **redes neuronales profundas** (Deep Learning) juegan un papel clave cuando se necesitan modelos capaces de detectar patrones extremadamente complejos en datos como imágenes, texto, audio o grandes volúmenes de información estructurada.

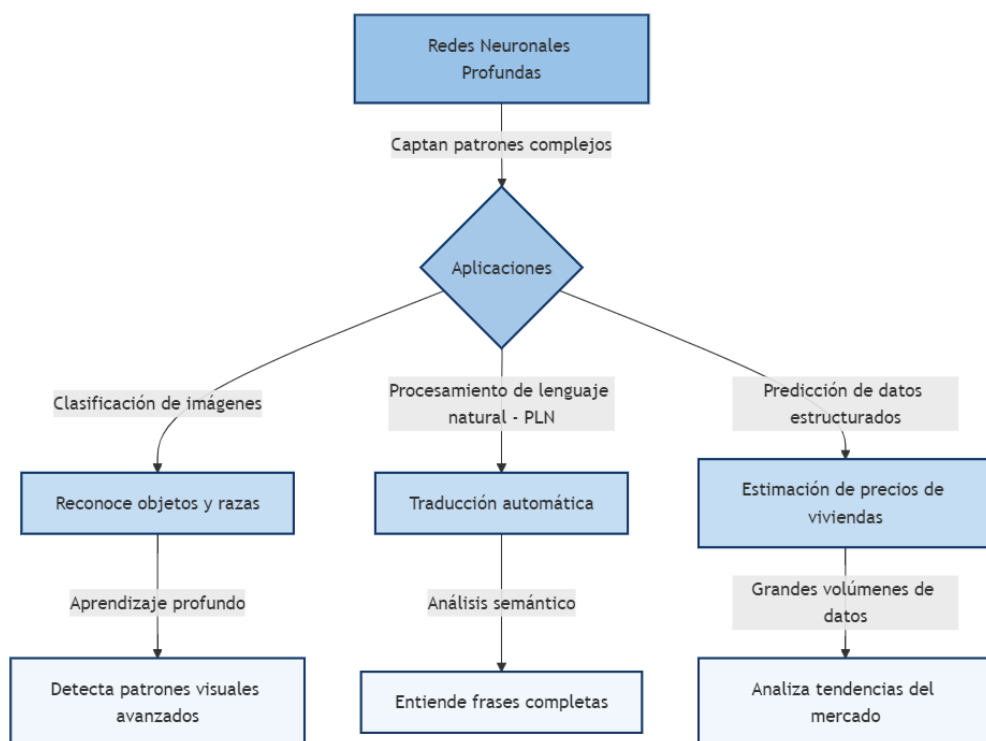
A diferencia de otros métodos más simples como la regresión lineal o los árboles de decisión, que funcionan bien con datos más estructurados y fáciles de interpretar, las redes neuronales profundas destacan en problemas donde la relación entre los datos de entrada y la salida no es evidente a simple vista. En este contexto, la complejidad del modelo se traduce en una mayor capacidad para captar detalles sutiles y combinaciones de características que otros enfoques no detectarían fácilmente.

Por ejemplo, en el caso de **clasificación de imágenes**, una red neuronal profunda no solo reconoce que una imagen contiene un gato o un perro, sino que también puede distinguir razas específicas basándose en patrones visuales avanzados. Lo hace gracias a su estructura en capas, donde las primeras identifican formas y bordes simples, y las más profundas combinan esa información para detectar patrones más abstractos.

En el ámbito del **procesamiento de lenguaje natural (PLN)**, como en sistemas de traducción automática o asistentes de voz, las redes neuronales profundas

permiten entender el significado de frases completas en lugar de analizar palabra por palabra de forma aislada. Esto es lo que hace posible que herramientas como Google Translate o ChatGPT generen respuestas más naturales y precisas.

Cuando se trata de **predicción de datos estructurados**, como en el ejemplo de estimar el precio de viviendas mencionado en el texto, las redes neuronales pueden ser útiles si se manejan grandes volúmenes de información con múltiples factores interconectados. Por ejemplo, podrían analizar no solo los datos básicos de una vivienda (tamaño, ubicación, número de habitaciones), sino también características más sutiles como la tendencia del mercado, la demanda en la zona y la evolución de precios en el tiempo.



Sin embargo, su uso no siempre es necesario ni recomendable. Si el problema se puede resolver con modelos más sencillos y explicables, como la regresión lineal o los árboles de decisión, es preferible optar por estos, ya que las redes neuronales profundas requieren más **datos, potencia de cálculo y tiempo de entrenamiento**. Además, su “caja negra” hace que interpretar sus decisiones sea más difícil, lo que puede ser un inconveniente en entornos donde se necesita justificar cada predicción.

2.2 ¿CÓMO ENTIENDE DEEPSEEK NUESTRAS PREGUNTAS?

Para comprender cómo DeepSeek interpreta nuestras preguntas, primero debemos imaginar el tipo de “inteligencia” que emplea para procesar el lenguaje humano. Como ya sabemos, a diferencia de un simple motor de búsqueda que identifica palabras clave y devuelve resultados exactos sin mayor análisis, DeepSeek va más allá. Utiliza modelos de inteligencia artificial entrenados específicamente para entender la intención de la pregunta, identificar los elementos relevantes que la componen y, a partir de ahí, ofrecer una respuesta precisa o una serie de resultados útiles. Este proceso implica descomponer la frase o el enunciado del usuario en sus partes más esenciales, relacionarlas con una base de conocimiento o una red de datos, y luego reconstruir una respuesta coherente.

2.2.1 Procesamiento del lenguaje natural (PLN) explicado

Cuando hacemos una pregunta a DeepSeek, el sistema no solo “ve” la palabra clave, sino que también trata de interpretar la relación semántica que hay entre las distintas palabras. Por ejemplo, si escribimos algo como “¿Cuál es la capital de Francia?”, un sistema de búsqueda tradicional se centraría en términos como “capital” y “Francia” para localizar una respuesta, mientras que DeepSeek, haciendo uso de algoritmos avanzados de comprensión lingüística, identificaría que el usuario solicita una ciudad específica asociada con un país. Esta diferencia sutil hace que la respuesta no sea simplemente un listado de páginas con la palabra “Francia” y “capital”, sino que se oriente directamente a la información solicitada: “París”.

Además, DeepSeek se adapta al contexto: si antes preguntaste algo sobre Europa o sobre el país galo, el sistema puede utilizar esa información previa para refinar la búsqueda. Así, se construye una especie de “memoria” a corto plazo que permite al motor responder de manera más ajustada y natural, acercándose a la forma en que interactuaría un ser humano en una conversación. Esta capacidad de análisis y adaptación proviene de la forma en que DeepSeek integra diversas técnicas de Procesamiento del Lenguaje Natural, reconocimiento de patrones y modelos de aprendizaje profundo.

El Procesamiento del **Lenguaje Natural (PLN)**, o NLP por sus siglas en inglés (Natural Language Processing), es la rama de la inteligencia artificial que se centra en permitir que las máquinas “entiendan” y “generen” lenguaje de manera similar a los seres humanos. El PLN se apoya en disciplinas como la lingüística, la informática y la estadística para analizar y extraer significado de textos o de discurso hablado. Una de sus bases principales es la segmentación de las oraciones en diferentes componentes lingüísticos, como palabras, sintagmas y oraciones, para así poder determinar qué función gramatical cumple cada elemento.



Para DeepSeek, el PLN resulta indispensable, ya que le permite diferenciar, por ejemplo, los sujetos y objetos de una frase, las relaciones temporales, y la estructura lógica que define la intención del usuario. Imaginemos que preguntas: “¿Qué tiempo hace hoy en Madrid?”. DeepSeek utilizará técnicas de PLN para descomponer la frase en tokens (palabras clave como “tiempo”, “hoy”, “Madrid”), identificar la intención principal (información meteorológica) y luego buscar en su base de datos o a través de servicios conectados cuál es la temperatura, la previsión de lluvia o la velocidad del viento.

Este proceso se vuelve complejo cuando nos enfrentamos a lenguaje ambiguo o con matices culturales y contextuales. Por ejemplo, la palabra “banco” puede referirse a una institución financiera o a un asiento. Para resolver estas ambigüedades, los modelos de PLN de DeepSeek tienen en cuenta el contexto que rodea a la palabra, así como el historial de búsquedas o el diálogo previo, de manera que identifiquen con la mayor precisión posible a qué se está refiriendo el usuario.

El PLN ha evolucionado de sistemas basados en reglas y gramáticas rígidas hacia arquitecturas de redes neuronales de aprendizaje profundo (Deep Learning). Como ya sabemos, estas redes se entrenan con enormes conjuntos de datos para “aprender” patrones estadísticos que ayudan a clasificar, agrupar y entender el texto de formas que anteriormente habrían requerido un análisis manual interminable. Para DeepSeek, esto se traduce en la capacidad de realizar inferencias más complejas y adaptarse mejor a cada estilo de comunicación, generando respuestas más naturales y precisas.

2.2.2 Reconocimiento de patrones y contexto

Además del PLN, un factor clave en la forma en que DeepSeek comprende nuestras preguntas es su capacidad de “reconocer patrones” tanto en el lenguaje como en la información previa. El reconocimiento de patrones implica que el sistema detecte regularidades, similitudes y estructuras repetitivas en grandes volúmenes de datos. Gracias a ello, DeepSeek puede anticipar qué tipo de información suele solicitarse cuando se plantea determinada pregunta o qué términos suelen aparecer asociados a cierto tema. Así, cuando alguien consulta “recomendaciones de restaurantes italianos en el centro”, el motor no solo ve las palabras “restaurantes” e “italianos”, sino que cruza datos de geolocalización, reseñas y horarios comerciales, encontrando patrones de calidad y popularidad que ayudarán a personalizar la respuesta.

El contexto es igualmente esencial. No siempre formulamos nuestras preguntas de manera clara y cerrada; a veces empezamos con dudas parciales o incompletas, o cambiamos repentinamente de un tema a otro. Mediante el uso de algoritmos de aprendizaje profundo, DeepSeek “recuerda” elementos de búsquedas anteriores o pistas dadas en mensajes previos para entender mejor el trasfondo de la conversación. Por ejemplo, si primero preguntas sobre “Museos famosos en París” y luego comentas “¿Cuánto cuesta la entrada?”, DeepSeek infiere que sigues hablando de los museos parisinos y no necesita que repitas explícitamente el nombre del lugar para saber a qué te refieres.



En la práctica, este reconocimiento de patrones y contexto se logra a través de modelos de lenguaje como las redes neuronales de tipo Transformer, que tienen la capacidad de procesar grandes secuencias de palabras y asignarles distintos “pesos” de relevancia para generar la respuesta más acertada. Esto se complementa con el análisis de la “historia” de interacciones en tiempo real, de modo que el sistema pueda hilar varias preguntas sucesivas y ofrezca información de manera coherente. Esto hace que la experiencia de usar DeepSeek sea más fluida y conversacional, alejándose de la idea de una simple búsqueda de palabras clave para acercarse a un diálogo inteligente.

A continuación, se presenta un ejemplo práctico que ilustra cómo DeepSeek emplea el reconocimiento de patrones y el análisis de contexto en una conversación real sobre planes en Madrid. En cada fase del diálogo, verás cómo el sistema retiene la información previa y adapta sus respuestas para ofrecer datos más pertinentes y coherentes.

Pregunta inicial y sugerencias generales

Usuario: *“Hola DeepSeek, ¿qué lugares de interés puedo visitar en Madrid este fin de semana?”*

Lugares de interés para visitar en Madrid

Hola DeepSeek, ¿qué lugares de interés puedo visitar en Madrid este fin de semana?”



¡Hola! Madrid es una ciudad llena de lugares increíbles para visitar. Aquí tienes algunas sugerencias para este fin de semana:

1. Museo del Prado

- Uno de los museos más importantes del mundo, con obras maestras de artistas como Velázquez, Goya y El Bosco.

2. Parque del Retiro

- Perfecto para un paseo relajante. Puedes alquilar una barca en el estanque, visitar el Palacio de Cristal o simplemente disfrutar de los jardines.

3. Palacio Real

- La residencia oficial de la Familia Real Española. No te pierdas los lujosos salones y la armería real.

4. Plaza Mayor y Puerta del Sol

- Dos de los lugares más emblemáticos de Madrid. Ideal para tomar un café o disfrutar de la animada vida callejera.

5. Museo Reina Sofía

- Aquí podrás ver el famoso cuadro "Guernica" de Picasso, junto con una impresionante colección de arte contemporáneo.

6. Mercado de San Miguel

- Un lugar perfecto para probar tapas y productos típicos españoles en un ambiente único.

7. Templo de Debod

- Un antiguo templo egipcio donado a España. Es un gran lugar para ver el atardecer.

Message DeepSeek

DeepThink (R1)

Search

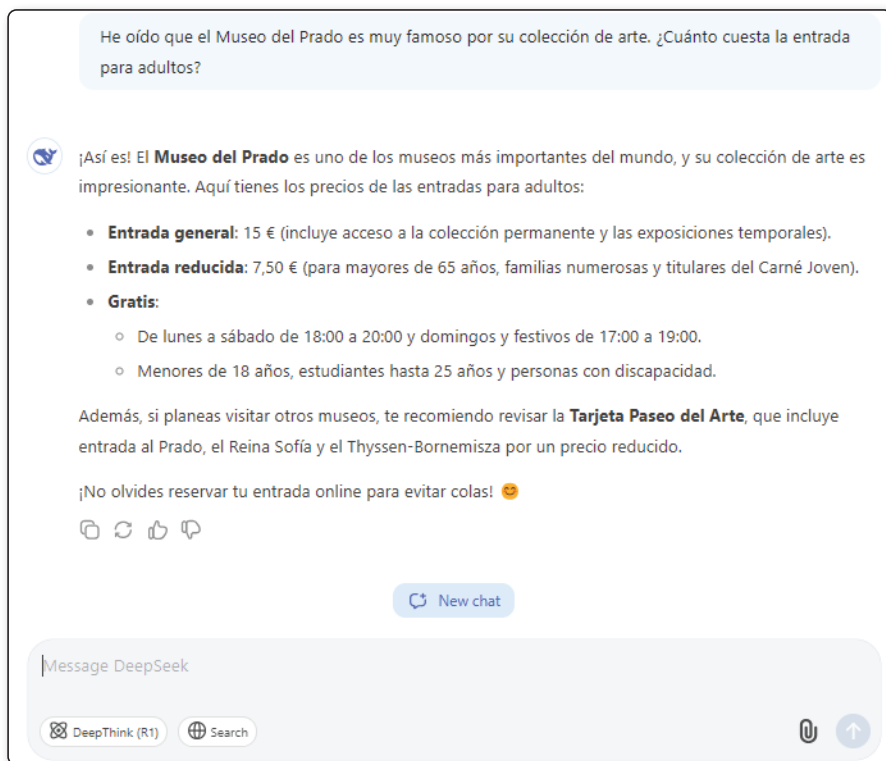


► Cómo actúa DeepSeek:

- En esta primera interacción, el sistema reconoce las palabras clave “lugares de interés”, “Madrid” y “este fin de semana”.
- A partir de estas pistas, DeepSeek consulta su base de datos o su modelo de conocimiento sobre qué sitios suelen ser relevantes para un plan turístico en la capital española, considerando también la cercanía temporal (“fin de semana”).
- El sistema no se limita a listar páginas web relacionadas con “Madrid”, sino que integra técnicas de Procesamiento de Lenguaje Natural (PLN) para comprender que el usuario quiere sugerencias de lugares para visitar, y no simplemente datos aleatorios.
- El resultado es una serie de recomendaciones bien estructuradas (Museo del Prado, Parque del Retiro, Palacio Real, etc.), cada una con una breve descripción.

Profundizando en un lugar específico:

Usuario: “He oído que el Museo del Prado es muy famoso por su colección de arte. ¿Cuánto cuesta la entrada para adultos?”



► **Cómo actúa DeepSeek:**

- Aquí se observa el **reconocimiento de patrones**: el usuario menciona “Museo del Prado” y plantea explícitamente una pregunta acerca del **precio de la entrada**.
- DeepSeek recuerda que hace un momento se habló de “lugares de interés en Madrid” y que dentro de la lista estaba precisamente el Museo del Prado.
- Mediante el **análisis de contexto**, el sistema asocia la frase “colección de arte” con un museo, confirmando que el usuario busca información sobre el coste de las entradas para adultos en ese lugar concreto.
- Los modelos de lenguaje tipo Transformer permiten que la IA tome en cuenta tanto las palabras usadas en la nueva consulta (“colección de arte”, “entrada para adultos”) como la referencia previa (el Museo del Prado en Madrid).
- La respuesta indica diferentes opciones de entrada (general, reducida, gratuita), reflejando cómo DeepSeek integra datos actualizados y específicos.

Ajuste de la respuesta según la nueva pregunta:

Usuario: “¿Hay algún día en que la entrada sea gratis y cómo reservo mis entradas?”

¿Hay algún día en que la entrada sea gratis y cómo reservo mis entradas?

 ¡Sí! El **Museo del Prado** ofrece **entrada gratuita** en horarios específicos, lo que es una gran oportunidad para visitarlo sin costo. Aquí te detallo cómo y cuándo:

Horarios de entrada gratuita:

- **De lunes a sábado:** de 18:00 a 20:00 horas.
- **Domingos y festivos:** de 17:00 a 19:00 horas.

Cómo reservar entradas:

- 1. Entrada gratuita:**
 - No es necesario reservar para acceder en los horarios gratuitos, pero te recomiendo llegar con anticipación, ya que suele haber colas.
- 2. Entrada de pago:**
 - Puedes reservar tus entradas online a través de la **página oficial del Museo del Prado:** museodelprado.es.
 - Selecciona la fecha y hora de tu visita.
 - Recibirás un código QR que deberás mostrar en la entrada.

Consejos:

- Si visitas en horario gratuito, intenta llegar al menos 30 minutos antes para evitar largas esperas.
- Si prefieres evitar aglomeraciones, la entrada de pago te permitirá acceder con mayor comodidad.

¡Disfruta de tu visita al Museo del Prado! 😊

 New chat

Message DeepSeek

 DeepThink (R1)  Search

▀ **Cómo actúa DeepSeek:**

- En este punto, el usuario desea información adicional: la **existencia de días con entrada gratuita** y el **método de reserva**.
- DeepSeek, otra vez, no pierde el hilo de la conversación. Gracias a la memoria de la interacción, sabe que se sigue hablando del Museo del Prado, aunque el usuario no lo mencione de forma explícita.
- El motor aplica el **PLN** para interpretar que “entrada gratuita” y “reservar mis entradas” siguen vinculados a la temática previa de las tarifas del museo, y no a otro sitio de interés.
- La respuesta se centra, por tanto, en ofrecer los datos concretos de los horarios con entrada gratuita (de lunes a sábado, domingos y festivos) y explica cómo adquirir las entradas, ya sean gratuitas o de pago.
- Gracias a los **modelos de lenguaje de tipo Transformer**, la IA puede “recordar” la conversación completa y relacionar estas nuevas preguntas con la información previa sin necesidad de que el usuario repita los detalles.

¿Por qué funciona?

1. Procesamiento de Lenguaje Natural (PLN):

- DeepSeek analiza la sintaxis y la semántica de la frase para extraer la intención. Las palabras clave y su posición en la frase son procesadas para generar la respuesta apropiada.
- En la pregunta sobre el precio de la entrada, el sistema entiende que la palabra “colección de arte” describe el tipo de museo; cuando luego pregunta por “cómo reservo mis boletos”, sabe que se refiere a la entrada al mismo lugar (Museo del Prado).

2. Reconocimiento de patrones:

- Mediante el entrenamiento con grandes volúmenes de datos, DeepSeek ha “aprendido” que las personas suelen preguntar precios, horarios, días gratuitos y métodos de reserva al buscar información sobre museos.
- Estos patrones le permiten generar respuestas completas y no solo centrarse en una parte de la duda.

3. Análisis de contexto y memoria conversacional:

- Cada paso de la conversación se “encadena” al anterior. Si el usuario omite volver a mencionar “Museo del Prado”, DeepSeek infiere que el tema no ha cambiado y puede responder con más detalle.
- Las redes neuronales de tipo Transformer utilizan mecanismos de auto-atención para determinar qué partes de la conversación previa siguen siendo relevantes en la pregunta actual.

4. Diálogo inteligente en lugar de búsqueda de palabras clave:

- A diferencia de un buscador tradicional que listaría sitios web con la expresión “precio entrada Museo del Prado”, DeepSeek construye una respuesta clara y directa.
- Esto se traduce en una experiencia más cercana a la interacción con un experto humano, capaz de razonar y recordar información de turnos de conversación anteriores.

Gracias a todo este proceso:

- El usuario obtiene una lista de recomendaciones específicas y detalladas sobre lugares de interés en Madrid.
- Puede profundizar en uno de ellos (Museo del Prado) pidiendo más información sobre precios, descuentos y horarios gratuitos, recibiendo datos claros sin tener que repetir el contexto.
- Conoce la posibilidad de reservar las entradas y hasta obtiene enlaces para consultar más detalles.
- El sistema se adapta a cada nueva consulta, demostrando una capacidad de mantener la coherencia a lo largo de varias preguntas sucesivas.

En conclusión, este ejemplo real evidencia cómo el **uso de redes neuronales de tipo Transformer** y el **PLN** permiten a DeepSeek reconocer patrones de conversación, analizar el contexto y ofrecer respuestas coherentes y útiles. Cada turno de la conversación está conectado al anterior, generando una interacción fluida y personalizada, muy diferente de las búsquedas estáticas basadas únicamente en palabras clave. De esta manera, DeepSeek logra acercarse más a la experiencia de dialogar con un experto humano en lugar de limitarse a devolver resultados desordenados e impersonales.

¿Cómo encaja todo esto en la arquitectura de redes neuronales tipo Transformer?

1. Procesamiento secuencial y auto-atención: las redes Transformer procesan el texto completo de forma que cada palabra (o token) “ve” a las demás dentro de la misma oración y a lo largo de la conversación previa. Esto permite que el sistema entienda que “Museo del Prado” sigue siendo relevante cuando el usuario pregunta por “la entrada” o “días gratuitos” en mensajes posteriores.
2. Asignación de pesos de relevancia: la técnica conocida como “self-attention” calcula la importancia de cada palabra en relación con las otras. En el ejemplo, “entrada” tiene un peso específico cuando se habla de precios y tickets; “Museo del Prado” retiene un alto nivel de relevancia mientras el usuario no cambie de tema.
3. Análisis de la “historia” de interacciones: DeepSeek no trata cada pregunta como un evento aislado. Más bien, analiza las conversaciones como un flujo en el que cada turno añade información sobre lo que el usuario desea. Si el usuario cambia abruptamente de tema a mitad de la conversación, la red ajusta ese cambio de contexto aprovechando los mismos mecanismos de auto-atención, pero modifica los pesos relativos a la temática anterior y se centra en la nueva.

Las principales ventajas de este enfoque son las siguientes:

- Mayor precisión: al mantener el contexto entre varias preguntas, DeepSeek responde con información más pertinente, evitando confusiones típicas de motores de búsqueda que se basan solo en palabras clave.
- Conversaciones naturales: el usuario no siente que deba repetir la misma información continuamente. DeepSeek recuerda la referencia a “Museo del Prado” sin necesidad de que se mencione en cada mensaje.
- Aprovechamiento de patrones: al haber visto múltiples variantes de la misma pregunta (“¿Cuánto vale la entrada?”, “¿Cuánto cuesta?”, “¿Cuál es el precio?”), el sistema reconoce la intención y unifica la respuesta, independientemente de la formulación exacta.
- Respuestas más completas: el modelo puede ofrecer datos adicionales (horarios, descuentos, etc.) basados en patrones previos de lo que habitualmente interesa a los usuarios.

2.3 COMPONENTES CLAVE DE DEEPSEEK

Para comprender la arquitectura completa de DeepSeek y por qué es tan eficaz en la búsqueda y procesamiento de información, es necesario detenerse en los componentes que lo hacen funcionar. De forma general, podemos decir que DeepSeek está construido sobre varios pilares que, trabajando en conjunto, ofrecen una experiencia de consulta casi conversacional. En primer lugar, destacan los módulos de búsqueda y procesamiento, que son el “motor” encargado de rastrear, extraer y analizar la información pertinente a cada consulta. A continuación, está el sistema de aprendizaje continuo, una pieza fundamental que garantiza la mejora constante del desempeño, adaptándose a los cambios en el lenguaje, en las fuentes de datos y en los intereses del usuario.



Cada uno de estos componentes cumple una misión específica, pero, al mismo tiempo, están interconectados para garantizar una interacción coherente y precisa. Por ejemplo, si un usuario realiza una pregunta sobre un tema complejo, los módulos de búsqueda y procesamiento se ponen en marcha para encontrar

información actualizada y filtrarla con precisión, mientras que el sistema de aprendizaje continuo incorpora esa experiencia a su acervo, extrayendo patrones que le permitirán responder con mayor acierto en futuras interacciones similares. La evolución de DeepSeek depende, en gran medida, de la capacidad de estos componentes de comunicarse entre sí, retroalimentarse y ajustarse a las necesidades cambiantes de los usuarios.

2.3.1 Módulos de búsqueda y procesamiento

Los módulos de búsqueda y procesamiento de DeepSeek constituyen el corazón operativo que permite, en primer lugar, **localizar** la información en una amplia base de datos o en la web, y luego **procesarla** para elaborar respuestas claras y ajustadas. Podríamos pensar en esta parte de la arquitectura como el “filtro inteligente” que, ante una pregunta, busca en un océano de datos solo aquellas referencias relevantes, las organiza y les da forma de respuesta.

1. **Indexación avanzada:**

Uno de los pasos clave en todo motor de búsqueda es la indexación. En el caso de DeepSeek, se lleva a cabo un proceso de indexación avanzado, donde además de almacenar palabras clave, se generan representaciones semánticas de los documentos. Esto implica analizar la relación entre diferentes términos y frases, de modo que el sistema entienda no solo que dos palabras aparecen juntas, sino también **cómo** y **por qué** se relacionan.

2. **Procesamiento semántico del lenguaje:**

Una vez localizados los documentos potencialmente relevantes, el sistema pasa a la fase de comprensión. Aquí, entra en juego el Procesamiento del Lenguaje Natural (PLN), que va mucho más allá de emparejar palabras clave. DeepSeek analiza la intención de la consulta, identifica sinónimos, reconoce construcciones gramaticales complejas e incluso comprende matices contextuales. Por ejemplo, si en la búsqueda se menciona “banco” en el contexto de un mobiliario urbano, DeepSeek diferenciará este significado de “banco” como entidad financiera, gracias a sus métodos de análisis semántico.

3. **Recuperación de información estructurada:**

Tras refinar y comprender la petición, el módulo de búsqueda extrae las partes más relevantes y las compila. Si la pregunta tiene una respuesta directa (por ejemplo, una fecha, un precio o una definición), DeepSeek la presentará. En caso de que se trate de un tema más amplio, el sistema

ordenará la información, destacando los puntos más importantes y ofreciendo enlaces o sugerencias adicionales para profundizar.

Estos módulos de búsqueda y procesamiento funcionan, por lo tanto, como un engranaje que combina velocidad y precisión, respaldado por algoritmos de machine learning entrenados para entender el lenguaje natural. A medida que aumenta el volumen de datos, estos algoritmos siguen optimizando sus procesos, asegurando que el tiempo de respuesta se mantenga bajo y la pertinencia de la información sea elevada.

Veamos un ejemplo para cada una de las tres fases principales: indexación avanzada, procesamiento semántico del lenguaje y recuperación de la información de forma estructurada.

1. Indexación avanzada

- Supongamos que DeepSeek recibe una gran colección de artículos científicos sobre biología marina. En lugar de limitarse a anotar palabras clave como “pez” o “océano”, el sistema lleva a cabo una indexación avanzada. Así, cuando un artículo menciona “ecosistema marino” y otro habla de “entorno submarino”, DeepSeek comprende que ambos describen el mismo fenómeno de manera distinta. Para lograrlo, el motor analiza no solo la frecuencia de aparición de las palabras, sino también la relación semántica entre ellas.
- Un usuario busca “¿Cuáles son los efectos del cambio climático en los arrecifes de coral?”. DeepSeek ya ha indexado cientos de documentos relacionados y sabe que términos como “acidez del océano”, “pérdida de biodiversidad” y “blanqueamiento del coral” están asociados entre sí. Gracias a esa indexación, puede “poner en la misma carpeta” estos temas, de modo que la búsqueda sea mucho más precisa y rápida.

2. Procesamiento semántico del lenguaje

- Una vez que DeepSeek identifica los documentos pertinentes, aplica técnicas de Procesamiento del Lenguaje Natural (PLN) para comprender la intención de la consulta y el contexto en el que se formulan las palabras.
- Si alguien pregunta “¿Cuál es el mejor banco en mi ciudad?”, DeepSeek analizará si se refiere a un banco financiero o a un banco para sentarse. Si detecta que la frase completa menciona “comisiones” y “horarios de apertura”, sabrá que la consulta se refiere a entidades financieras. Por el contrario, si la pregunta se enfoca en “dónde sentarse” o “parques

públicos”, el sistema inferirá que se trata de un mobiliario urbano. Esto se debe a que DeepSeek no se queda solo con la palabra “banco”, sino que estudia la oración completa y su contexto.

3. Recuperación de información estructurada

- En la fase final, DeepSeek ordena los datos extraídos y los presenta al usuario de forma clara y útil.
- Imagina que preguntas “¿Cuál es el horario de entrada del Museo del Prado en Madrid?”. Si existe una respuesta concreta (por ejemplo, “de 10:00 a 20:00 horas”), DeepSeek la mostrará directamente. Pero si tu pregunta es más amplia, como “¿Qué museos puedo visitar en España?”, el sistema listará opciones, agrupará por ciudad o temática y añadirá enlaces a fuentes oficiales para que puedas profundizar. En lugar de ofrecer un bloque de texto desordenado, DeepSeek mostrará secciones o puntos clave de la información, facilitándote la navegación.

2.3.2 Sistema de aprendizaje continuo

Si los módulos de búsqueda y procesamiento son el músculo de DeepSeek, el **sistema de aprendizaje continuo** es, sin duda, su cerebro en constante evolución. Este componente se basa en la idea de que un asistente o motor de búsqueda no puede quedarse estancado con el conocimiento inicial con el que se entrenó, sino que debe **actualizarse, refinarse y expandirse** de manera perpetua.

1. Aprendizaje supervisado y no supervisado:

Para que DeepSeek sea capaz de absorber nuevas tendencias en el lenguaje o en los datos, se aplican distintas técnicas de aprendizaje. En el aprendizaje supervisado, los desarrolladores y expertos le proporcionan ejemplos concretos, indicándole cómo debe clasificar o responder a ciertos tipos de preguntas. Por otro lado, el aprendizaje no supervisado le permite al sistema descubrir patrones por sí mismo en grandes volúmenes de datos, sin intervención humana directa. Es en este equilibrio entre supervisión y autonomía donde DeepSeek logra una flexibilidad extraordinaria.

2. Actualización de modelos en tiempo real:

A través de un sistema de retroalimentación, DeepSeek examina cómo los usuarios reaccionan a las respuestas que da: si encuentran la información útil, si plantean preguntas de seguimiento o si reformulan la búsqueda.

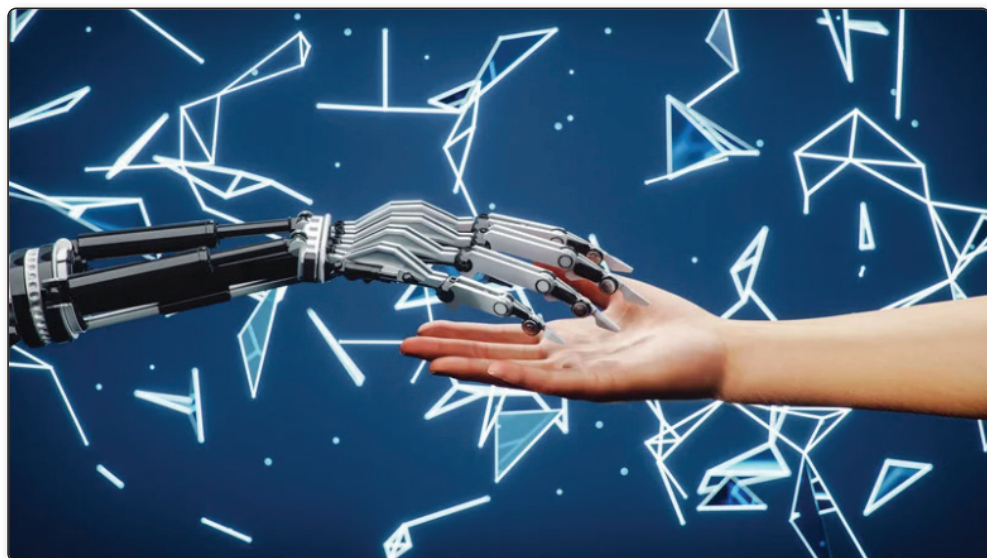
Todo esto se almacena y se “digieren” estadísticamente, de modo que el motor de IA aprenda de cada interacción. De esta forma, cuando un usuario plantea una pregunta similar en el futuro, el sistema no empezará desde cero, sino que aplicará las lecciones aprendidas de casos pasados.

3. **Detección de nuevas fuentes y actualización de contenido:**

El conocimiento humano se expande a gran velocidad, y la información puede quedar obsoleta en poco tiempo. Por ello, el sistema de aprendizaje continuo de DeepSeek está diseñado para rastrear nuevas fuentes de información (páginas web, artículos científicos, noticias, bases de datos especializadas, etc.) y evaluar su fiabilidad. Con cada actualización, se integran estos datos en la base de conocimientos, garantizando que las respuestas estén al día y reflejen los cambios que ocurren en el mundo real.

4. **Personalización inteligente:**

Además de aprender de manera global, DeepSeek puede adaptarse a los intereses particulares de cada usuario (siempre respetando su privacidad y las normativas de protección de datos). A medida que una persona utiliza el motor de búsqueda y hace más preguntas, DeepSeek registra patrones de preferencia, por ejemplo, sobre el estilo de información preferido (más técnico, más sencillo), la temática de mayor interés o la frecuencia con que se requieren actualizaciones. Todo esto se traduce en un sistema de recomendaciones más acertado y ajustado a la medida.



Veamos ejemplos específicos de cómo funciona cada una de las partes de este sistema de aprendizaje.

1. Aprendizaje supervisado y no supervisado

- DeepSeek combina ambas modalidades para ser más versátil.
- Imagina que los desarrolladores quieren que DeepSeek detecte recetas veganas con mayor precisión. En un enfoque de aprendizaje supervisado, se le proporcionan numerosos ejemplos etiquetados de recetas veganas y no veganas, indicándole a la IA dónde se equivoca. Poco a poco, DeepSeek “aprende” a reconocer patrones (ingredientes típicos, sustitutos de la carne, etc.).
- En aprendizaje no supervisado, el sistema recibe un gran volumen de recetas sin etiquetas. DeepSeek busca por sí mismo similitudes y diferencias, agrupa recetas que comparten características y descubre tendencias sin que nadie se lo indique de forma explícita. Luego, integra esas relaciones con lo aprendido en los ejemplos supervisados.

2. Actualización de modelos en tiempo real

- Uno de los puntos clave de DeepSeek es que no se limita a “adivinar” la respuesta inicial y olvidarla. Sigue aprendiendo de cada interacción y ajusta sus modelos internos.
- Si un usuario pregunta repetidamente sobre un tema de física cuántica y califica la respuesta como poco precisa, DeepSeek registra esa retroalimentación. El motor analiza la razón por la que la respuesta se inadecuó, revisa otras fuentes o ejemplos y corrige el modelo para la siguiente ocasión. De esta forma, la próxima vez que alguien formule una pregunta similar, DeepSeek ya habrá mejorado su respuesta basándose en la experiencia anterior.

3. Detección de nuevas fuentes y actualización de contenido

- El conocimiento humano y las noticias del mundo real se actualizan de manera constante. DeepSeek rastrea continuamente fuentes de alta calidad para incorporar datos recientes.
- Aparece un nuevo estudio científico que modifica lo que se sabía sobre la eficacia de una vacuna. DeepSeek analiza la publicación, determina su fiabilidad (por ejemplo, si proviene de una revista indexada y revisada por pares) y actualiza su base de conocimientos. Así, cuando un usuario pregunte sobre los resultados de ese estudio, DeepSeek estará listo para ofrecer la información más reciente.