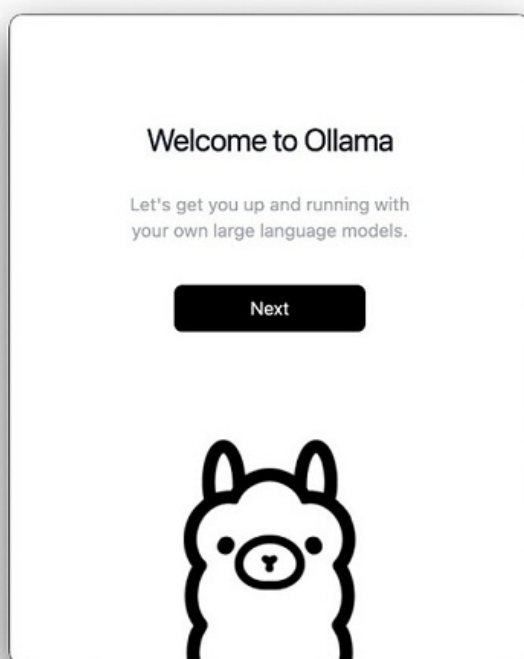


Qué es Ollama y cómo usarlo para instalar en tu ordenador modelos de inteligencia artificial como Llama, DeepSeek y más

Este programa te permite instalar modelos de inteligencia artificial en tu ordenador y usarlos de forma local sin conectarte a Internet.

Vamos a explicarte **qué es Ollama y cómo funciona** esta aplicación con la que puedes instalar DeepSeek en tu ordenador, así como otros modelos de inteligencia artificial como **Llama, Phi, Mistral, Qwen, Llava o Gemma**. Estos modelos quedarán instalados de forma local para que no tengas que ir a ninguna web para usarlos, y para que todo lo que hagas se quede en tu ordenador.



Vamos a empezar el artículo explicándote qué es exactamente Ollama, y todo lo que necesitas saber sobre este programa. Luego, pasaremos a decirte cómo instalarlo en tu ordenador, y luego cómo empezar a usarlo para instalar un modelo de inteligencia artificial y usarlo.

Qué es y cómo funciona Ollama

Ollama es un programa que puedes instalar en cualquier ordenador, tanto con sistema operativo Windows como con macOS o GNU/Linux. Se trata de un cliente de modelos de inteligencia artificial, por lo que **es la base sobre la que luego instalar una IA** que quieras utilizar.

Ollama tiene dos particularidades. La primera es que **te permite usar una IA de forma local**. Esto quiere decir que en vez de ir a la página de chat con inteligencia artificial de una empresa, el modelo está instalado en tu ordenador y lo utilizas directamente sin entrar en ninguna web.

Esto te favorece de tres maneras. Primero porque **los datos de todo lo que haces se quedan en tu PC**, de forma que ninguna empresa los utiliza. Segundo porque **puedes usar la IA sin conexión** a internet, y tercero porque puedes saltarte censuras que tenga un modelo de inteligencia artificial que estés utilizando en una web. Eso sí, lo que no puedes es hacer búsquedas por internet.

Y la segunda particularidad es que **funciona a través de la terminal de tu ordenador**, o el símbolo de sistema en Windows. Esto hace que no tengas que usar una aplicación aparte. Cuando instales Ollama, luego tendrás que usar la consola de tu dispositivo para instalar y ejecutar en ella el modelo que quieras, y las preguntas y los prompts los escribes en la consola, donde también tendrás las respuestas.

Cómo instalar Ollama



Get up and running with large language models.

Run Llama 3.3, DeepSeek-R1, Phi-4, Mistral, Gemma 2, and other models, locally.

Download ↓

Available for macOS,
Linux, and Windows

Para empezar a usar Ollama, lo primero que tienes que hacer es **instalar su aplicación para ordenador**. Para eso entra en su web ollama.com, y pulsa en el botón *Download* que te aparecerá.

Download Ollama



macOS



Linux



Windows

Download for Windows

Requires Windows 10 or later

Ahora, irás a la página donde tienes que **elegir el sistema operativo** para el que quieres bajarte el programa. Una vez lo hayas elegido, pulsa en el botón *Download*. Por defecto la web mostrará el sistema que estás usando, pero podrás descargar el ejecutable de cualquier otro.

Cuando lo descargues, lanza el programa de instalación. Instalar Ollama es muy sencillo, solo tienes que pulsar en el botón de siguiente en la pantalla de presentación, y luego **pulsar en el botón *Install*** en la pantalla de instalación.

Cómo usar Ollama

Q

Search models

All

Embedding

Vision

Tools

Popular

▼

deepseek-r1

DeepSeek's first-generation of reasoning models with comparable performance to OpenAI-o1, including six dense models distilled from DeepSeek-R1 based on Llama and Qwen.

1.5b 7b 8b 14b 32b 70b 671b

↓ 3M Pulls 🏷 28 Tags ⌚ Updated 7 days ago

llama3.3

New state of the art 70B model. Llama 3.3 70B offers similar performance compared to the Llama 3.1 405B model.

tools 70b

↓ 910.9K Pulls 🏷 14 Tags ⌚ Updated 7 weeks ago

phi4

Phi-4 is a 14B parameter, state-of-the-art open model from Microsoft.

14b

↓ 209.5K Pulls 🏷 5 Tags ⌚ Updated 2 weeks ago

Una vez has instalado Ollama, lanza la aplicación. Verás que no pasa nada, esto es porque tienes que **abrir el terminal de tu ordenador**, que en Windows se llama símbolo de sistema. Ahora, antes de empezar tienes que ir a la web donde verás **todos los modelos de IA disponibles**. La web es ollama.com/search.

deepseek-r1

DeepSeek's first-generation of reasoning models with comparable performance to OpenAI-o1, including six dense models distilled from DeepSeek-R1 based on Llama and Qwen.

1.5b 7b 8b 14b 32b 70b 671b

3.1M Pulls Updated 7 days ago

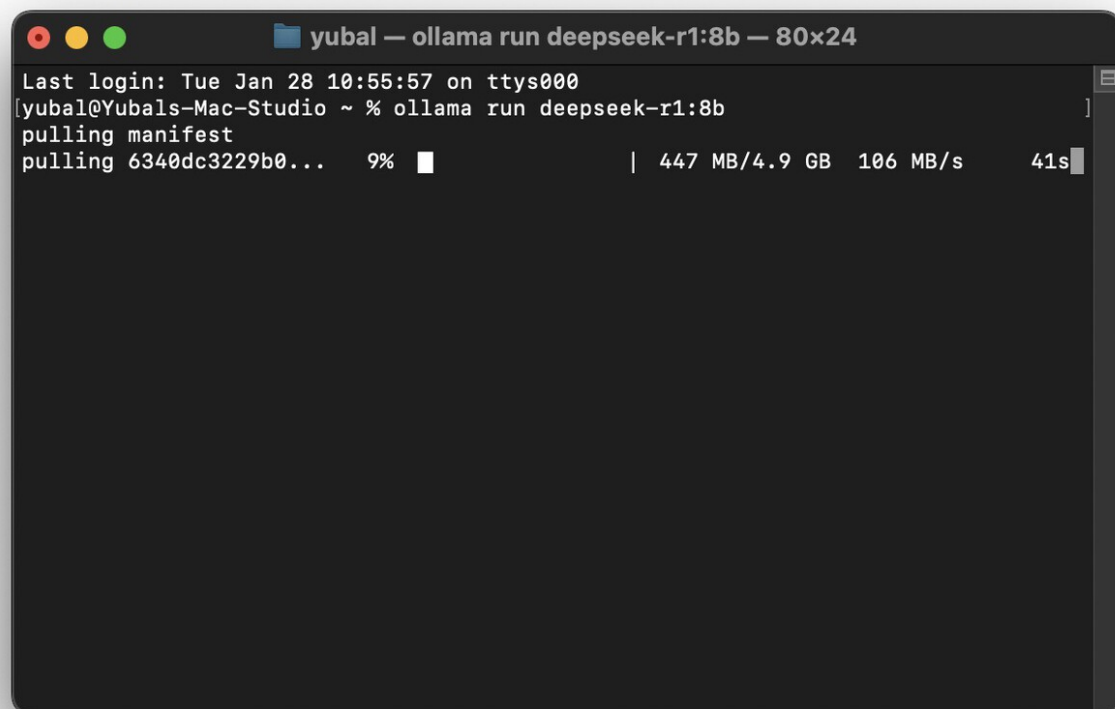
8b

28 Tags

ollama run deepseek-r1:8b

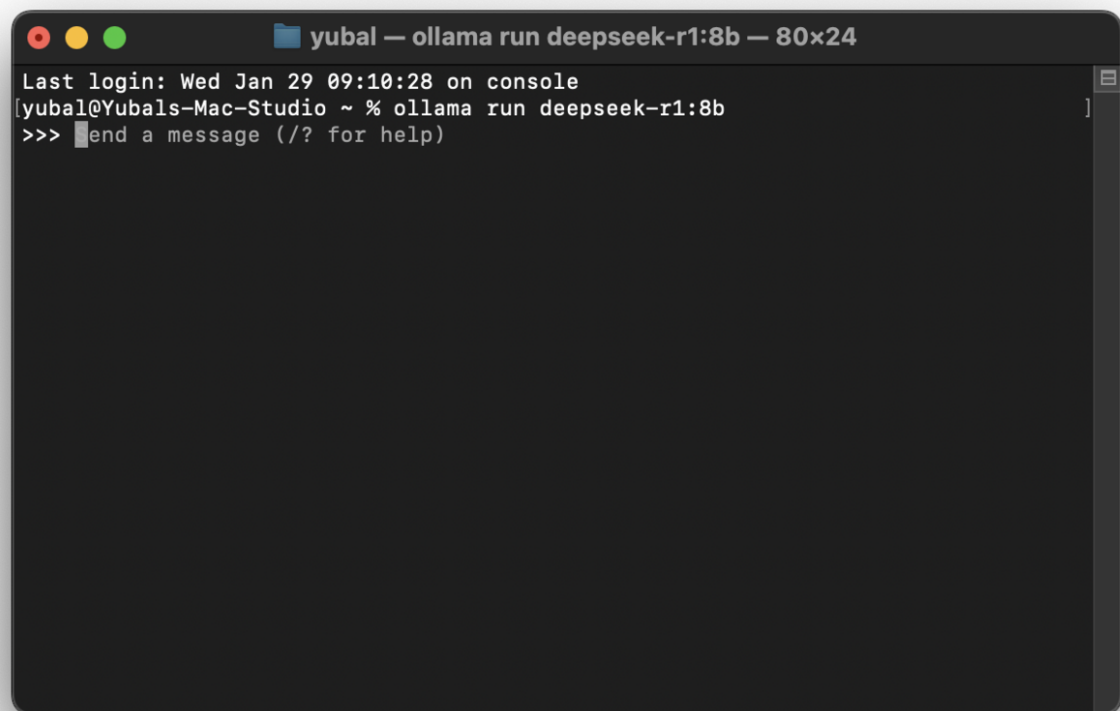
Updated 8 days ago		28f8fd6cdc67 · 4.9GB
model	arch llama · parameters 8.03B · quantization Q4_K_M	4.9GB
params	{ "stop": ["< begin_of_sentence >", "< end_of_sentence >..."	148B
template	{{- if .System }}{{ .System }}{{ end }} {{- range \$i, \$_ := ...	387B
license	MIT License Copyright (c) 2023 DeepSeek Permission is hereb...	1.1kB

Pulsa en uno de los modelos de la lista para entrar en su ficha. Aquí tendrás toda la información, y también tendrás un selector para elegir las distintas versiones disponibles. Cuando elijas una versión, arriba a la derecha tienes el comando necesario para instalarlo y lanzarlo en Ollama.

A terminal window titled 'yubal — ollama run deepseek-r1:8b — 80x24'. The window shows the following text: 'Last login: Tue Jan 28 10:55:57 on ttys000', '[yubal@Yubals-Mac-Studio ~ % ollama run deepseek-r1:8b]', 'pulling manifest', and 'pulling 6340dc3229b0... 9% [progress bar] | 447 MB/4.9 GB 106 MB/s 41s'. The progress bar is a small black rectangle. The terminal has a dark background and standard macOS window controls at the top.

```
yubal — ollama run deepseek-r1:8b — 80x24
Last login: Tue Jan 28 10:55:57 on ttys000
[yubal@Yubals-Mac-Studio ~ % ollama run deepseek-r1:8b
pulling manifest
pulling 6340dc3229b0... 9% [progress bar] | 447 MB/4.9 GB 106 MB/s 41s
```

Ahora, solo tienes que **usar el comando para lanzar el modelo en tu terminal**, como por ejemplo `ollama run deepseek-r1:8b` para lanzar la versión 8b de [DeepSeek](#) R1. La primera vez que uses el comando primero se instalará el modelo, pero las próximas ya lo lanzarás directamente. Recuerda que para hacer esto en la terminal de tu ordenador primero tienes que haber lanzado la aplicación de Ollama.

A terminal window with a dark background and light gray text. The title bar at the top reads 'yubal — ollama run deepseek-r1:8b — 80x24'. The terminal content shows a login message: 'Last login: Wed Jan 29 09:10:28 on console'. Below that is the prompt '[yubal@Yubals-Mac-Studio ~ % ollama run deepseek-r1:8b]'. The next line shows the prompt '>>>' followed by a cursor and the text 'end a message (/? for help)'.

```
yubal — ollama run deepseek-r1:8b — 80x24
Last login: Wed Jan 29 09:10:28 on console
[yubal@Yubals-Mac-Studio ~ % ollama run deepseek-r1:8b ]
>>> end a message (/? for help)
```

Y tras escribir el comando, **se lanzará el modelo de IA en el terminal**. Esto lo vas a distinguir porque ves que en el campo de escritura del terminal ahora aparece un >>>, lo que significa que lo que escribas se lo enviarás al modelo de inteligencia artificial.

```
yubal — ollama run deepseek-r1:8b — 76x20
[pulling 0cb05c6e4e02... 100% 487 B]
verifying sha256 digest
writing manifest
success
>>> Qué es Xataka?
<think>

</think>

Xataka es un sitio web de noticias y contenido general, fundado en 2004,
que cubre temas como tecnología, ciencia, política, economía y cultura. Es
conocido por ser una fuente de información actualizada y diversificada,
además de incluir secciones como noticiosas, ifestyle, opiniones y más. El
sitio está disponible en español y otras lenguas, y es bastante popular en
[Several países hispanoamericanos.

Si estás buscando algo específico o necesitas más información, déjame
saber! 😊
```

Ahora, en la línea de comandos de tu ordenador podrás **escribir el prompt que quieras** lanzarle a la IA que hayas elegido, y tras unos segundos empezará a generarte la respuesta.

Fuente: xataka.