

Trabajo Final - Tutorial: Instalación de n8n

Integrantes del trabajo: Julián Derudi | Leandro Tittarelli | Lucas Lencina.

Objetivo del trabajo

El objetivo es documentar paso a paso cómo:

- Instalar n8n (herramienta de automatización de flujos)
- Instalar Ollama (motor para ejecutar modelos LLM localmente)
- Descargar y correr un modelo de lenguaje liviano
- Superar errores comunes que pueden surgir en el proceso

Requisitos previos

- Sistema operativo: Ubuntu 20.04+ (probado en Ubuntu Desktop)
- Acceso a terminal con permisos de administrador (sudo)
- Conexión a internet
- Al menos 4 GB de RAM para modelos básicos (más recomendable: 8 GB)
- Espacio en disco: mínimo 10 GB libres

Parte 1 – Instalación de Node.js, npm y n8n

1 Actualizar el sistema

```
sudo apt update && sudo apt upgrade -y
```

2 Instalar Node.js (versión LTS recomendada)

```
curl -fsSL https://deb.nodesource.com/setup_lts.x | sudo -E bash -  
sudo apt install -y nodejs
```

Verificar:

```
node -v  
npm -v
```

3 Instalar n8n de forma global

```
npm install -g n8n
```

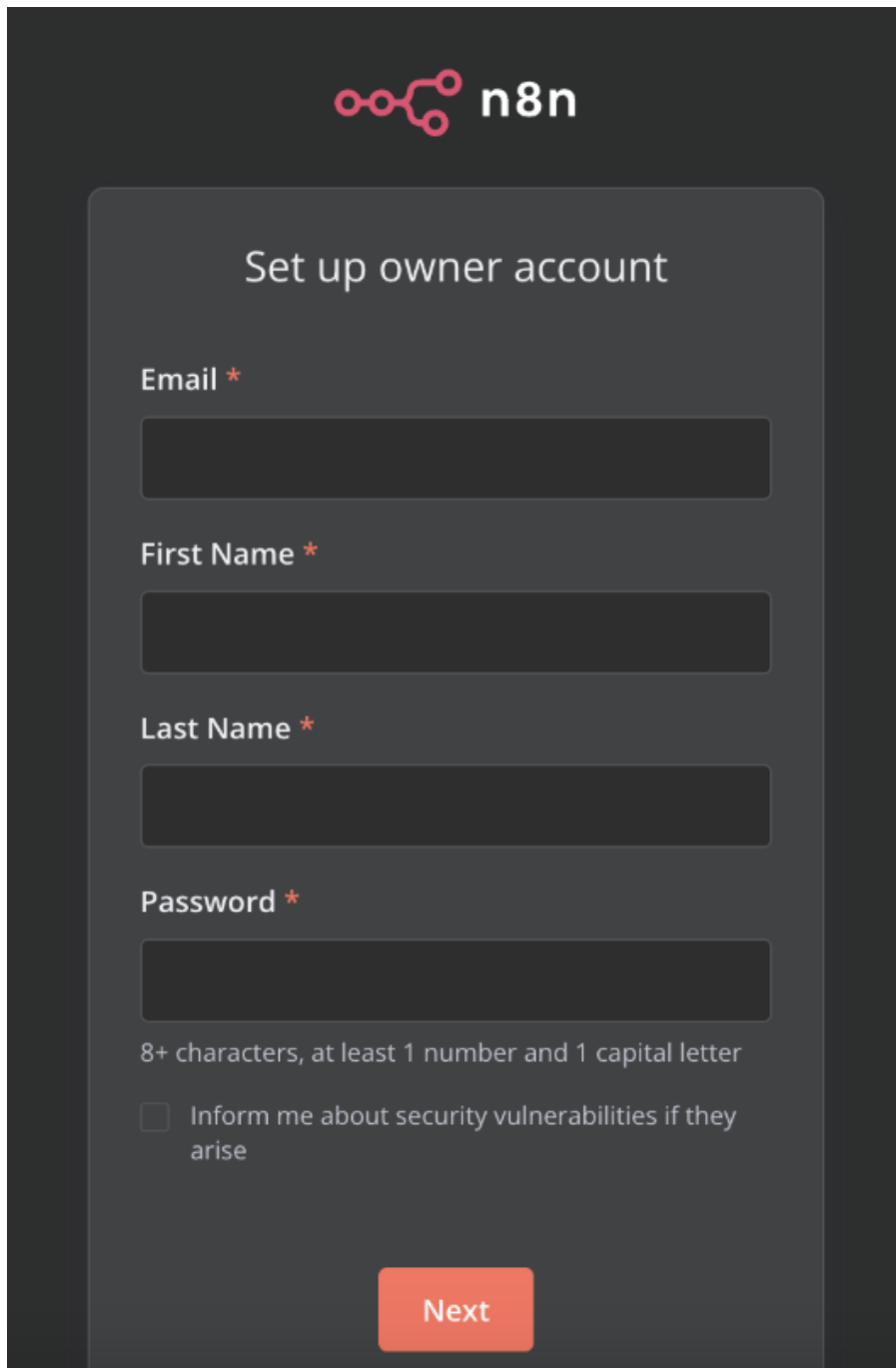
4 Ejecutar n8n

n8n

Se accede desde el navegador en: <http://localhost:5678>

Al ejecutar n8n por primera vez y acceder a <http://localhost:5678>, te aparecerá un formulario para crear la cuenta de administrador. Esto configura el acceso básico para proteger tu instancia de n8n.

Luego de esto, n8n guarda los datos en su base de datos (por defecto una base SQLite local si no configuraste otra). Quedará protegida con usuario y contraseña.



The screenshot shows the n8n logo at the top, followed by the title 'Set up owner account'. Below this are four input fields: 'Email *', 'First Name *', 'Last Name *', and 'Password *'. Each field is represented by a dark gray rectangular box. Below the password field, there is a text requirement: '8+ characters, at least 1 number and 1 capital letter'. At the bottom of the form, there is a checkbox labeled 'Inform me about security vulnerabilities if they arise' and a red 'Next' button.

Set up owner account

Email *

First Name *

Last Name *

Password *

8+ characters, at least 1 number and 1 capital letter

☐ Inform me about security vulnerabilities if they arise

Next

Posibles errores y soluciones con n8n

Error: EACCES: permission denied

Causa: No hay permisos para escribir en /usr/lib/node_modules

Solución: Usar sudo: `$ sudo npm install -g n8n`

Error: Command 'n8n' not found

Causa: La instalación falló o no se añadió al PATH

Solución: Verificar PATH y reinstalar con sudo

Parte 2 – Instalación de Ollama

1 Descargar e instalar Ollama

```
curl -fsSL https://ollama.com/install.sh | sh
```

Advertencia común al instalar Ollama

WARNING: No NVIDIA/AMD GPU detected. Ollama will run in CPU-only mode.

Esto no es un error, solo una advertencia: la ejecución será más lenta porque usa el procesador.

Parte 3 – Descargar y ejecutar un modelo de lenguaje

✗ Primer intento fallido: mistral

```
ollama run mistral
```

Falla con PCs de poca RAM:

Error: model requires more system memory (5.5 GiB) than is available (3.7 GiB)

✓ Solución: usar modelo liviano gemma:2b

```
ollama run gemma:2b
```

Alternativa: agregar SWAP

```
sudo fallocate -l 4G /swapfile  
sudo chmod 600 /swapfile  
sudo mkswap /swapfile  
sudo swapon /swapfile
```

Para dejarlo permanente:

```
echo '/swapfile none swap sw 0 0' | sudo tee -a /etc/fstab
```

Estado final esperado

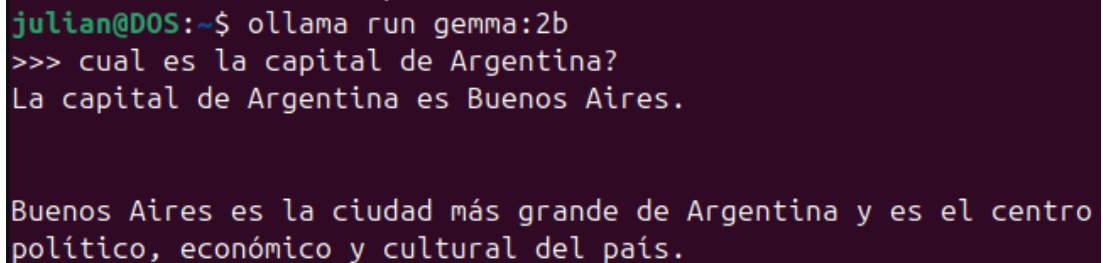
n8n: Abrir <http://localhost:5678>

ollama: Ejecutar `ollama run gemma:2b`

RAM: Verificar con `free -h` que haya suficiente disponible

Ejemplo para verificar que anda (desde la terminal)

Desde la terminal, siguiendo los pasos deberías de ver una respuesta parecida a esta:



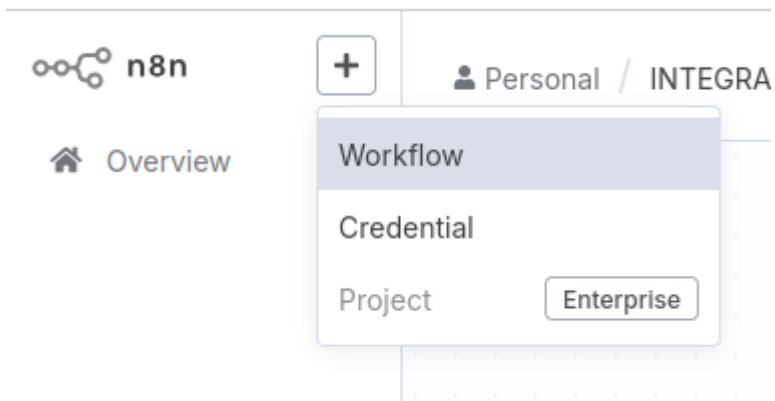
```
julian@DOS:~$ ollama run gemma:2b  
>>> cual es la capital de Argentina?  
La capital de Argentina es Buenos Aires.  
  
Buenos Aires es la ciudad más grande de Argentina y es el centro  
político, económico y cultural del país.
```

Flujo demostrativo

Ahora vamos a ingresar a n8n desde la consola y vamos a crear un workflow

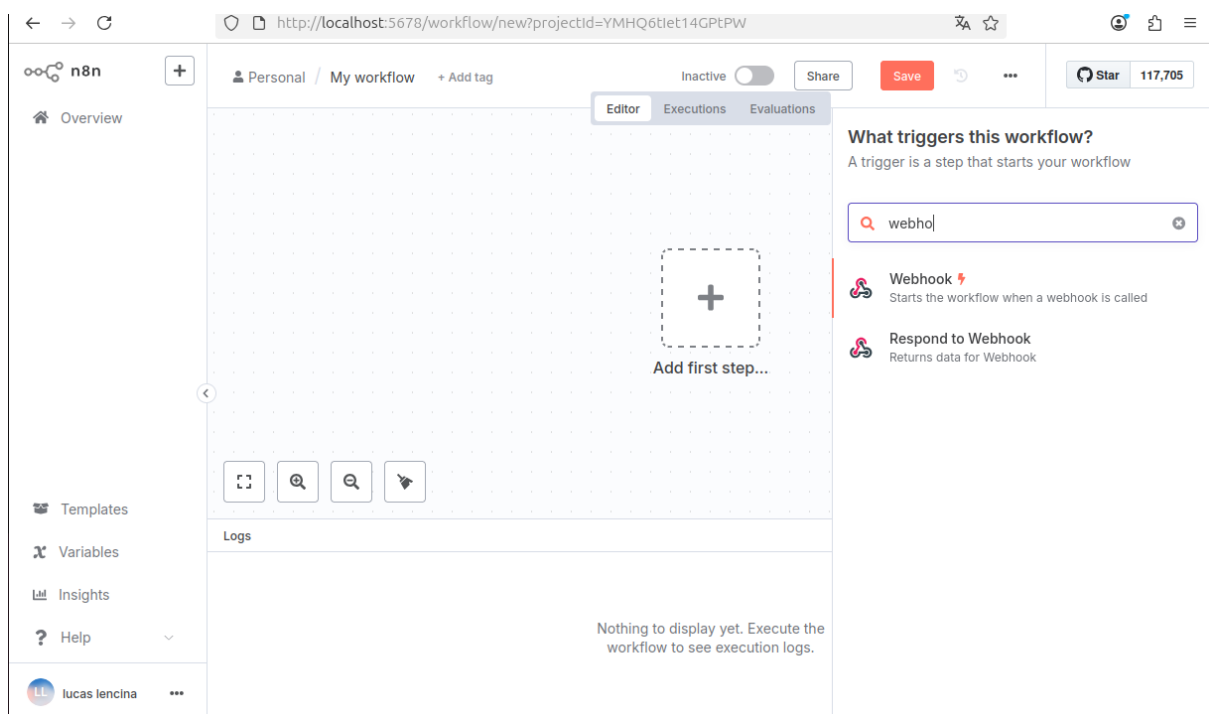
luego presionamos la letra “O” eso nos abra por defecto n8n desde el navegador y ingresamos con el usuario y contraseña que creamos anteriormente.

A continuacion creamos un workflow presionando el + junto a n8n como podemos ver en la siguiente imagen, despues seleccionamos workflow lo cual nos llevara a crear un workflow vacio en cual debemos cambiarle el nombre y guardarlo.

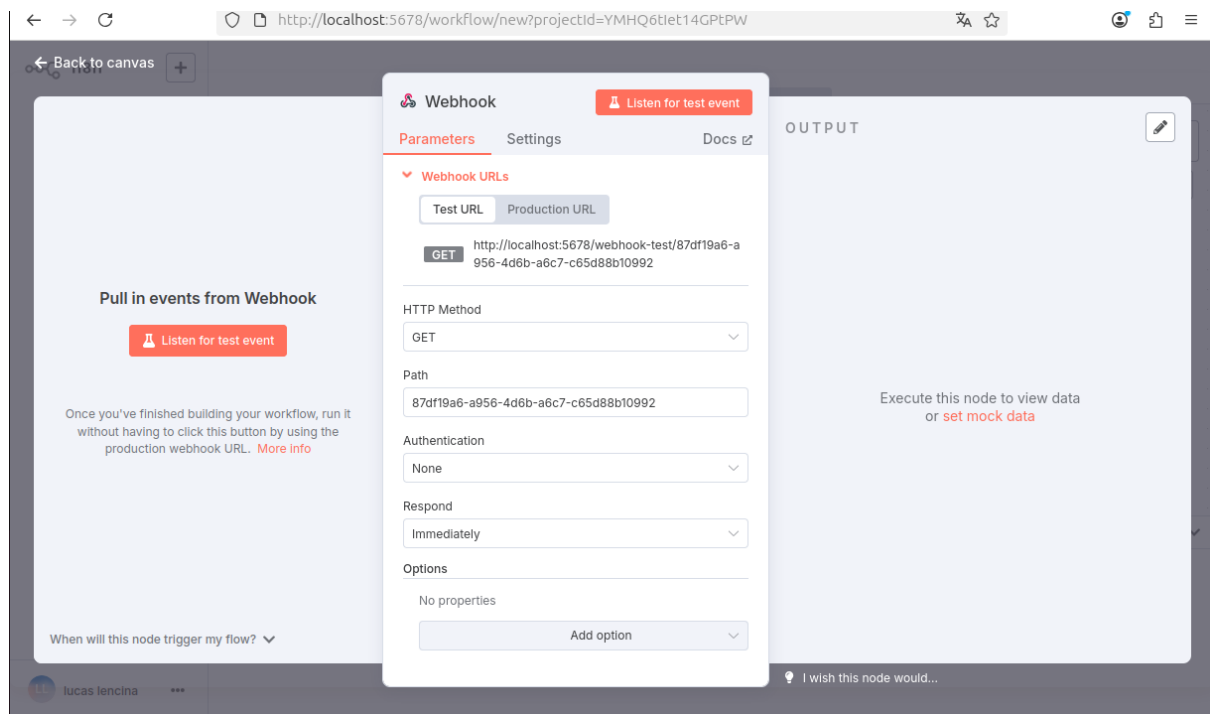


```
pache98@pache98: ~  
his is ignored for now, but in the future n8n will attempt to change the permissions automatically. To automatically enforce correct permissions now set N8N_ENFORCE_SETTINGS_FILE_PERMISSIONS=true (recommended), or turn this check off set N8N_ENFORCE_SETTINGS_FILE_PERMISSIONS=false.  
User settings loaded from: /home/pache98/.n8n/config  
Initializing n8n process  
n8n ready on ::, port 5678  
  
There is a deprecation related to your environment variables. Please take the recommended actions to update your configuration:  
- N8N_RUNNERS_ENABLED -> Running n8n without task runners is deprecated. Task runners will be turned on by default in a future version. Please set `N8N_RUNNERS_ENABLED=true` to enable task runners now and avoid potential issues in the future. Learn more: https://docs.n8n.io/hosting/configuration/task-runners/  
  
[license SDK] Skipping renewal on init because renewal is not due yet or cert is not initialized  
Version: 1.97.1  
  
Editor is now accessible via:  
http://localhost:5678  
  
Press "o" to open in Browser.
```

para el siguiente paso nos vamos a dirigir a la esquina superior derecha del workflow donde se encuentra el signo + eso nos abra un menu de busqueda donde debemos buscar el nodo webhook y agregarlo

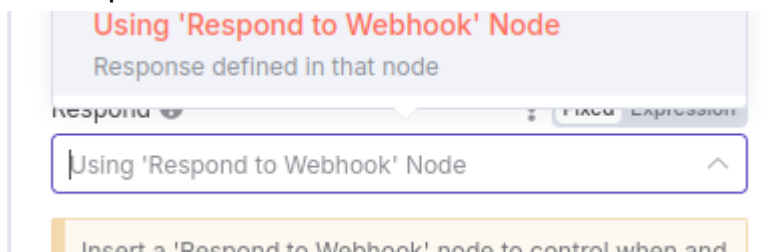


una vez agregado el nodo debemos configurarlo



Configuracion del nodo:

1. En HTTP Method seleccionamos la opcion POST
2. Path preguntar (<https://localhost:5678/webhook/preguntar>)
3. En Respond seleccionamos



Luego guardamos el workflow y activamos el workflow.

Conclusiones

- Instalar y usar n8n y Ollama localmente es posible incluso sin GPU ni mucha RAM
- Es clave tener en cuenta los requisitos de memoria al elegir el modelo LLM
- La combinación de n8n + Ollama permite crear flujos potentes, 100% offline
- Este tutorial cubre no solo los pasos ideales, sino los errores reales y sus soluciones