

NIMLAS Workshop 2026 Manual: Installing Ollama and Running Your First Gemma Model Using the Ollama App

Workshop Goal

By the end of this exercise, you will be able to:

- Install Ollama on your computer.
- Download a small open-source Gemma language model.
- Interact with the model using the Ollama desktop application.
- Verify that everything is working correctly.

No programming experience is required.

Why Are We Using the Ollama App Instead of Command Prompt?

There are two common ways to use Ollama.

Table 1: Comparison of Ollama App Use

Command Prompt / Terminal	Ollama Desktop App
You type text commands such as <code>ollama run gemma3:1b</code> .	You click buttons and type prompts in a graphical interface.
Preferred by programmers and advanced users.	Preferred by beginners and users who are not comfortable with programming.
Makes automation and scripting easier.	Easier to learn and use.

Both methods use the same language models and produce the same answers.

The only difference is **how you interact with the model**.

Before You Begin

Operating Systems Supported

This guide supports:

- Microsoft Windows 10 or 11
- macOS (Apple Silicon M1/M2/M3 or Intel)

Internet Connection

You will need an Internet connection during installation because the model must be downloaded.

After the model has been downloaded, it can be used without an Internet connection.

Disk Space

Reserve at least **5 GB** of free disk space.

Note: If you prefer, you can watch this YouTube video for visual instruction.

Link to the video: <https://youtu.be/7LEvSOiTWZk?si=8BXtQH1phf029EX->

Part 1. Install Ollama

Windows

Step 1

Open your web browser and visit

Link to the website: <https://ollama.com>

Step 2

Click **Download**.

Choose **Windows**.

Step 3

Open the downloaded installer.

If Windows asks

"Do you want to allow this app to make changes to your device?"

click **Yes**.

Step 4

Follow the installation instructions.

The default settings are appropriate for most users.

Step 5

Click **Finish** when installation is complete.

macOS

Step 1

Open your web browser and visit

Link to the website: <https://ollama.com>

Step 2

Click **Download**.

Choose **macOS**.

Step 3

Open the downloaded installer.

Step 4

Drag **Ollama** into the **Applications** folder.

Step 5

Open **Applications** and launch **Ollama**.

If macOS asks whether you want to open software downloaded from the Internet, click **Open**.

Part 2. Launch the Ollama Application

After installation,

start the Ollama application.

On Windows,

you can find it from the **Start Menu**.

On macOS,

you can find it in the **Applications** folder.

When Ollama is running, you may also see its icon in the system tray (Windows) or menu bar (macOS).

Part 3. Download the Gemma Model

Inside the Ollama application,

select **Models** (or the model library, depending on your version).

Search for

Gemma

Choose the smallest available Gemma model.

For this test run, use

Gemma 3 1B

Click **Download**.

The download may take several minutes depending on your Internet connection.

You only need to download the model once.

Part 4. Start a New Chat

After the download finishes,

click **New Chat**.

Choose

Gemma 3 1B

from the available models.

You should now see a message box where you can type questions.

Part 5. Ask Your First Question

Type:

Hello! Please introduce yourself in one paragraph.

Press **Enter**.

The model should respond after a few seconds.

Congratulations!

You are now running an LLM entirely on your own computer.

Part 6. Try a Few More Questions

Example 1

Summarize the following paragraph in two sentences.

Paste one of the interview excerpts from today's workshop.

Example 2

Identify the three main themes discussed in this interview.

Example 3

Extract the participant's main challenges and return the results as a table.

Example 4

Explain the difference between qualitative and quantitative research in language that a high school student can understand.

Part 7. Start Another Conversation

If you would like to begin a completely new conversation,
click

New Chat

and select the same model again.

Starting a new chat prevents earlier questions from influencing later answers.

Part 8. View Installed Models

Open the **Models** page.

You should see

Gemma 3 1B

listed as one of your installed models.

This confirms that the model has been downloaded successfully.

Part 9. Remove a Model (Optional)

If you need to free disk space later,

open the **Models** page.

Select the model.

Choose **Delete** or **Remove**.

You can download it again at any time.

Tips for Better Results

As you learned earlier in today's workshop, good prompts produce better answers.

Instead of asking

Summarize this.

try asking

Summarize this interview in three sentences. Identify the participant's positive and negative experiences separately. Use only information explicitly stated in the interview.

Small improvements in your prompt can greatly improve the quality of the output.

Troubleshooting

Problem 1: I cannot find the Ollama application.

Windows

Open the Start Menu and search for

Ollama

macOS

Open the Applications folder and look for

Ollama

Problem 2: The model is still downloading.

Downloads occur only once.

Please wait until the download finishes before starting a chat.

Problem 3: I cannot find the Gemma model.

Refresh the Models page.

If it still does not appear,

search for

Gemma

again.

Problem 4: The model responds very slowly.

This usually means your computer is busy or the model is relatively large.

Close unnecessary applications and try again.

Problem 5: My computer reports insufficient memory.

Close other programs and restart Ollama.

If the problem continues, try a smaller model available on Ollama.

Problem 6: I accidentally closed Ollama.

Simply reopen the application.

Your downloaded models will still be available.

You do not need to download them again.

Problem 7: The application appears to freeze.

Wait one or two minutes if the model has just started.

The first response is often slower because the model is being loaded into memory.

Congratulations!

You have successfully:

- Installed Ollama.
- Downloaded an open-source Gemma model.
- Started a local language model using a graphical interface.
- Asked your first questions.
- Learned how to begin new conversations and continue exploring on your own.

You are now ready to redo the hands-on exercises in today's workshop using a locally running LLM.