

The Local AI Stack: Building a Sovereign Machine Learning Workstation with Hyper-V, WSL2, and GPU Virtualization

Table of Contents

1. Research hardware requirements including SLAT-capable CPUs and NVIDIA Tensor Core GPUs for AI workloads.
2. Enable hardware virtualization (VT-x/AMD-V) and IOMMU support within the BIOS/UEFI settings.
3. Activate the Hyper-V feature set via the Windows Features dialog or PowerShell `Enable-WindowsOptionalFeature`.
4. Install the Virtual Machine Platform and Windows Subsystem for Linux (WSL) optional components.
5. Configure the system to use WSL 2 as the default version for all future Linux distribution installs.
6. Download and install the latest Linux kernel update package for Windows from Microsoft.
7. Install a high-performance Linux distribution like Ubuntu 22.04 LTS from the Microsoft Store or via CLI.
8. Set up the `.wslconfig` file in the user profile to manage global resource limits for all WSL 2 instances.
9. Allocate specific CPU core counts and memory limits within `.wslconfig` to prevent host OS starvation.
10. Implement custom swap file locations on high-speed NVMe storage to improve virtualization paging.
11. Install the NVIDIA Windows GPU Driver (Game Ready or Studio) to support GPU Paravirtualization.
12. Download and install the NVIDIA Container Toolkit to allow Docker to interface with the GPU in WSL 2.
13. Install Docker Desktop and configure it to use the WSL 2 backend for improved performance.
14. Alternatively, install Podman or Rancher Desktop for an open-source containerization approach.
15. Configure the NVIDIA runtime in the Docker `daemon.json` file to ensure container access to CUDA.
16. Verify GPU acceleration inside WSL 2 using the `nvidia-smi` command to check for driver connectivity.
17. Install the CUDA Toolkit inside the Linux distribution, ensuring version compatibility with host drivers.
18. Configure cuDNN libraries within the virtualized Linux environment for optimized deep learning.
19. Optimize VHDX storage by moving the WSL 2 virtual disk to a dedicated non-system partition.
20. Use the `wsl --mount` command to attach physical disks directly to the Linux environment for raw I/O.
21. Configure Hyper-V Virtual Switches (Internal vs. External) to manage network isolation for AI labs.
22. Implement Discrete Device Assignment (DDA) for passing through an entire GPU to a Hyper-V VM.

23. Explore GPU Partitioning (GPU-PV) scripts to share a single GPU across multiple Hyper-V VMs.
24. Setup Miniconda or Mamba within the WSL 2 environment for isolated Python environment management.
25. Install PyTorch and TensorFlow with specific CUDA support flags to leverage hardware acceleration.
26. Configure VS Code with the Remote - WSL extension for a seamless Windows-to-Linux development flow.
27. Set up Jupyter Lab servers within WSL 2 and map ports to access them via the Windows browser.
28. Enable WSLg (WSL GUI) to run Linux-native visualization tools like TensorBoard or Matplotlib windows.
29. Implement Nested Virtualization to allow running Hyper-V or Docker containers within a Hyper-V VM.
30. Adjust Windows Defender exclusions for WSL 2 project folders to improve file system performance.
31. Configure SSH keys for secure remote access to the Windows-centric AI virtualization stack.
32. Use PowerShell to automate the creation and teardown of Hyper-V VMs for ephemeral AI training tasks.
33. Setup a local model repository using Ollama or LM Studio within the virtualized layer.
34. Integrate Hugging Face CLI tools for efficient model weight management across environments.
35. Configure Git LFS in the Linux layer to handle large AI model files and datasets.
36. Optimize network throughput using the Mirrored network mode in WSL 2 (experimental features).
37. Setup a bridge network for Hyper-V VMs to allow them to act as independent nodes on the local network.
38. Monitor system performance using Windows Performance Monitor (PerfMon) and Linux htop/nvtop.
39. Implement automated backups of WSL 2 distributions using the wsl --export command.
40. Create a recovery plan for VHDX corruption or host system failures.
41. Fine-tune the Windows kernel power management settings to 'High Performance' for consistent AI throughput.
42. Test the stack by running a sample LLM (Large Language Model) inference benchmark.
43. Document the specific versions of CUDA, cuDNN, and drivers used to ensure stack reproducibility.
44. Evaluate the use of Windows Sandbox for isolated testing of untrusted AI libraries or scripts.
45. Finalize the production-ready AI stack by hardening network ports and user permissions within the Windows Firewall.

Example:

Step 17

Install the CUDA Toolkit inside the Linux distribution, ensuring version compatibility with host drivers. The Mathematical Engine – Deploying the CUDA Toolkit. In the previous steps, we established the communication bridge between the Windows host and the Linux guest. We verified that `nvidia-smi` functions correctly, proving that the vGPU is active. However, `nvidia-smi` only confirms the presence of the driver. To actually execute AI workloads—to compile kernels, run PyTorch operations, or manage tensor memory—the Linux distribution requires the CUDA Toolkit.

In a virtualized WSL2 environment, installing the CUDA Toolkit requires a surgical approach. A standard installation often attempts to bundle a Linux driver, which will catastrophically break the `dxgkrnl` bridge we have built. This step details the precision installation of the toolkit while ensuring absolute compatibility with the host Windows driver.

The CUDA Compatibility Matrix

Before downloading a single package, you must understand the CUDA Forward Compatibility model as it applies to virtualization.

The Host Ceiling: Your Windows NVIDIA Driver defines the maximum CUDA version your hardware can support.

The Guest Toolkit: You can install a CUDA Toolkit inside WSL2 that is equal to or lower than the version supported by the host driver.

WSL2 Specifics: Unlike native Linux, WSL2 supports "CUDA Minor Version Compatibility," meaning a newer toolkit can often run on an older driver, provided the driver meets a specific minimum baseline (typically WDDM 3.0+).

The Engineering Rule: Always check the output of `nvidia-smi` on the Windows host. If it says "CUDA Version: 12.4", your Linux toolkit should ideally be avoiding the "Driver Overwrite" Trap.

The most common failure point in AI virtualization is the `apt-get install cuda` command. By default, this meta-package includes the NVIDIA Linux driver as a dependency. If this driver installs, it will overwrite the specialized WSL2 mapping, rendering the GPU invisible.

The Solution: The cuda-toolkit Meta-package

We must install the `cuda-toolkit-X-Y` package instead of the naked `cuda` package. This installs the compilers (`nvcc`), headers, and libraries without touching the driver layer.

Installation: The Targeted WSL2 Repository

Inside your Ubuntu 22.04 LTS instance, we will use the specialized NVIDIA repository. This repository contains builds specifically flagged for WSL2 compatibility.

Step 1: Remove the Outdated GPG Key

```
sudo apt-key del 7fa2af80
```

Step 2: Setup the Pinning Logic

To ensure that Ubuntu's standard package manager doesn't try to "update" your CUDA libraries with generic versions, we pin the NVIDIA repository.

```
wget  
https://developer.download.nvidia.com/compute/cuda/repos/wsl-ubuntu/x86_64/cuda-wsl-ubuntu.pin  
sudo mv cuda-wsl-ubuntu.pin /etc/apt/preferences.d/cuda-repository-pin-600
```

Step 3: Fetch the Toolkit

Replace 12-4 with the version that matches your host driver.

```
wget  
https://developer.download.nvidia.com/compute/cuda/repos/wsl-ubuntu/x86_64/cuda-keyring_1.1-1_all.deb  
sudo dpkg -i cuda-keyring_1.1-1_all.deb  
sudo apt-get update  
sudo apt-get -y install cuda-toolkit-12-4
```

Configuring the Environment Fabric

The toolkit is installed in `/usr/local/cuda-X.Y`, but the Linux shell does not yet know where the binaries are located. We must update the system paths to enable the compiler and libraries.

Modify the `.bashrc` or `.zshrc`:

```
export PATH=/usr/local/cuda-12.4/bin${PATH:+:${PATH}}  
export  
LD_LIBRARY_PATH=/usr/local/cuda-12.4/lib64${LD_LIBRARY_PATH:+:${LD_LIBRARY_PATH}}
```

Apply the changes:

```
source ~/.bashrc
```

The CUDA "Stub" and Library Linking

In a virtualized environment, the actual CUDA driver library (`libcuda.so`) is provided by the Windows host and mapped into `/usr/lib/wsl/lib`. The CUDA Toolkit provides a "stub" library for compilation.

Advanced Verification:

Verify that your system is correctly linking to the virtualized driver rather than a local file:

```
ldconfig -p | grep libcuda.so
```

Expected Result:

The path should point to `/usr/lib/wsl/lib/libcuda.so.1`. If it points to `/usr/local/cuda/...`, your AI frameworks will fail to initialize the GPU because they are trying to use a static library instead of the hypervisor's projected driver.

Post-Install Verification: Compiling a Test Kernel

The final test of the intelligence layer is compiling a raw CUDA kernel. This confirms that the headers, the compiler (nvcc), and the driver bridge are all in sync.

Step 1: Check the Compiler

```
nvcc --version
```

Step 2: Run a Simple Matrix Multiplication (Optional but recommended)

If you have the CUDA samples installed:

```
cd /usr/local/cuda/samples/1_Uutilities/deviceQuery
sudo make
./deviceQuery
```

Result: If deviceQuery returns Result = PASS, your virtualized AI stack is officially "Compute Ready."

The Acceleration Layer – cuDNN and NCCL With the CUDA Toolkit active, you have the base mathematical operations. However, for Deep Learning (CNNs, Transformers, etc.), you require cuDNN (CUDA Deep Neural Network library) and NCCL (NVIDIA Collective Communications Library) for multi-GPU scaling.

Installing cuDNN for WSL2

cuDNN provides highly tuned implementations for standard routines such as forward and backward convolution, pooling, normalization, and activation layers.

Download: Fetch the cuDNN local installer for Linux (Ubuntu 22.04 x86_64) from the NVIDIA Developer portal.

Integration:

```
sudo dpkg -i cudnn-local-repo-ubuntu2204-8.x.x.x_1.0-1_amd64.deb
sudo cp /var/cudnn-local-repo-*/cudnn-local-*.keyring.gpg /usr/share/keyrings/
```

```
sudo apt-get update
sudo apt-get install libcudnn8 libcudnn8-dev libcudnn8-samples
```

NCCL: Scaling Across the Hypervisor

If your research led you to a dual-GPU setup (e.g., two RTX 4090s), you must install NCCL. NCCL allows the two GPUs to communicate via the PCIe fabric (or NVLink) with minimal CPU intervention.

```
sudo apt-get install libnccl2 libnccl-dev
```

Ebook Conclusion: The Fully Realized Intelligence Rig

You have transitioned from a standard Windows installation to a masterfully orchestrated AI Virtualization Layer.

The Stack Summary:

Silicon: SLAT-capable CPU & Tensor Core GPUs.

Firmware: IOMMU, VT-x, and Re-size BAR enabled.

Hypervisor: Windows Server 2025 or Windows 11 running as a Type-1 Hypervisor.

Environment: WSL2 with a custom Microsoft Linux kernel.

Resource Control: `.wslconfig` prevents host starvation and optimizes NVMe paging.

Compute: NVIDIA Windows Driver bridged to a Linux CUDA Toolkit 12.x.

Orchestration: Docker/Podman/Rancher with the NVIDIA Container Runtime.

The Capability:

You can now pull the latest Llama 3, Stable Diffusion XL, or custom PyTorch models, knowing that the underlying virtualization layer is robust, isolated, and performing at 97-99% of bare-metal speeds.

By separating the "Intelligence Layer" from the "Host OS," you have created a workstation that is easy to backup, easy to scale, and virtually impossible to break through standard software conflicts. You are now ready to build the future of AI.

Step 18

Configure cuDNN libraries within the virtualized Linux environment for optimized deep learning. The Optimization Primitives – Integrating cuDNN for Deep Learning

With the CUDA Toolkit providing the base mathematical architecture, the environment is capable of general-purpose GPU computing. However, for modern AI

workloads—specifically those involving Convolutional Neural Networks (CNNs) and Transformers—raw CUDA is not enough. To achieve peak TFLOPS, we must layer in the NVIDIA CUDA Deep Neural Network library (cuDNN).

In a virtualized WSL2 environment, cuDNN acts as a tuning fork. It provides highly optimized implementations for standard routines such as forward and backward convolution, pooling, normalization, and activation layers. Without cuDNN, frameworks like PyTorch and TensorFlow would resort to generic CUDA kernels, resulting in a 30% to 50% performance penalty in training throughput.

The cuDNN Architecture in a Virtualized Context

cuDNN is a "primitives" library. It doesn't run code itself; rather, it provides a library of pre-optimized "engines" that deep learning frameworks call upon.

Heuristic Selection: cuDNN uses an internal heuristic to choose the fastest algorithm for a specific GPU's architecture (e.g., choosing a different convolution strategy for an RTX 4090 vs. an A100).

Virtualization Transparency: Because cuDNN sits entirely in "User Space" and communicates with the GPU via the CUDA API, it is largely unaware that it is running inside a Hyper-V partition. However, it relies heavily on the pinned memory and DMA (Direct Memory Access) paths we optimized.

Version Parity: The CUDA/cuDNN Handshake

The most common cause of "Runtime Error: cuDNN failed to initialize" is a version mismatch. Since we installed CUDA 12.4 in the previous steps, we must select a cuDNN version (typically v8.9 or v9.x) that is explicitly compiled for the 12.x toolkit.

cuDNN Version	Supported CUDA	Optimization Focus
v8.9.x 11.x - 12.x	Ampere/Ada	Stability
v9.1.x (Latest) 12.x	Blackwell	FP4/FP6 support, Graph API

Installation Method A: The Debian Local Repository (Recommended)

For our Ubuntu 22.04 LTS environment, the package manager approach is preferred because it handles the dependencies—like Zlib, which cuDNN uses for compressing descriptor data.

Step 1: Download the local installer

Visit the NVIDIA Developer portal and download the .deb for Ubuntu22.04 x86_64. Step 2: Install the repository GPG key

```
sudo dpkg -i cudnn-local-repo-ubuntu2204-8.9.7.29_1.0-1_amd64.deb
# Copy the keyring to the system folder
```

```
sudo cp /var/cudnn-local-repo-ubuntu2204-8.9.7.29/cudnn-local-*-keyring.gpg
/usr/share/keyrings/
```

Step 3: Deploy the Libraries

```
sudo apt-get update
sudo apt-get install libcudnn8 libcudnn8-dev libcudnn8-samples
```

Note: Installing libcudnn8-dev is mandatory for AI engineers, as it includes the header files (cudnn.h) required to compile custom C++ extensions for PyTorch.

Installation Method B: Manual Library Injection (The Researcher's Choice)

If you are managing multiple AI projects that require different cuDNN versions, you may prefer to install the libraries manually into the CUDA folder. This prevents system-wide conflicts.

Extract the Tarball:

```
tar -xvf cudnn-linux-x86_64-8.x.x.x_cuda12-archive.tar.xz
```

Copy to the Virtualized CUDA Path:

```
sudo cp cudnn-*-archive/include/cudnn*.h /usr/local/cuda/include
sudo cp -P cudnn-*-archive/lib/libcudnn* /usr/local/cuda/lib64
sudo chmod a+r /usr/local/cuda/include/cudnn*.h /usr/local/cuda/lib64/libcudnn*
```

Verification: Testing the Neural Handshake

Verification in a virtualized environment is critical to ensure the LD_LIBRARY_PATH we set is correctly picking up the cuDNN symbols.

The cuDNN Sample Test:

```
# Copy samples to a writable directory
cp -r /usr/src/cudnn_samples_v8/ ~/
cd ~/cudnn_samples_v8/mnistCUDNN
make clean && make
./mnistCUDNN
```

Expected Result:

If the virtualization bridge, the CUDA Toolkit, and the cuDNN libraries are all aligned, the output will conclude with:

Test passed!

Troubleshooting "Library Not Found":

If the test fails, run `ldconfig -p | grep cudnn`. If it returns empty, the Linux dynamic linker hasn't indexed the new libraries. Run:

```
sudo ldconfig
```

Optimizing cuDNN Performance in WSL2

To squeeze the final 5-10% of performance out of your virtualized intelligence layer, you must tune how cuDNN manages its memory workspace.

cuDNN Benchmarking: In your PyTorch code, always enable the benchmark flag. This instructs cuDNN to run a "warm-up" where it tests multiple algorithms on your vGPU and selects the fastest one for your specific hardware.

```
torch.backends.cudnn.benchmark = True
```

Environment Variable Tuning: For debugging or forcing specific precision, use the following variables in your `.bashrc`:

```
export CUDNN_LOGDEST_DBG=stdout: View real-time optimization logs.
```

```
export CUDNN_LOGLEVEL_DBG=1: Useful for spotting "Unsupported Algorithm" errors that occur when VRAM is fragmented.
```

Summary of the Deep Learning Fabric

By configuring cuDNN within the virtualized Linux environment, you have completed the High-Performance AI Stack.

Your Windows-centric setup is now functionally indistinguishable from a bare-metal Linux AI server, with the added benefits of Windows-based data management, UI accessibility, and Hyper-V's robust isolation. You are now prepared to deploy LLMs, train custom vision models, and lead advanced AI research from a single, unified virtualized workspace.

[THE END]