

H100 Tensor Core Architecture: A New Paradigm in Computation

Author: BenchChem Technical Support Team. **Date:** April 2026

Compound of Interest

Compound Name: H100

Cat. No.: B15585327

[Get Quote](#)

The **H100** introduces several architectural innovations that deliver substantial performance gains over its predecessors.[1] Built on a custom TSMC 4N process, the **GH100** GPU packs 80 billion transistors, enabling significant enhancements in processing power and efficiency.[2][3]

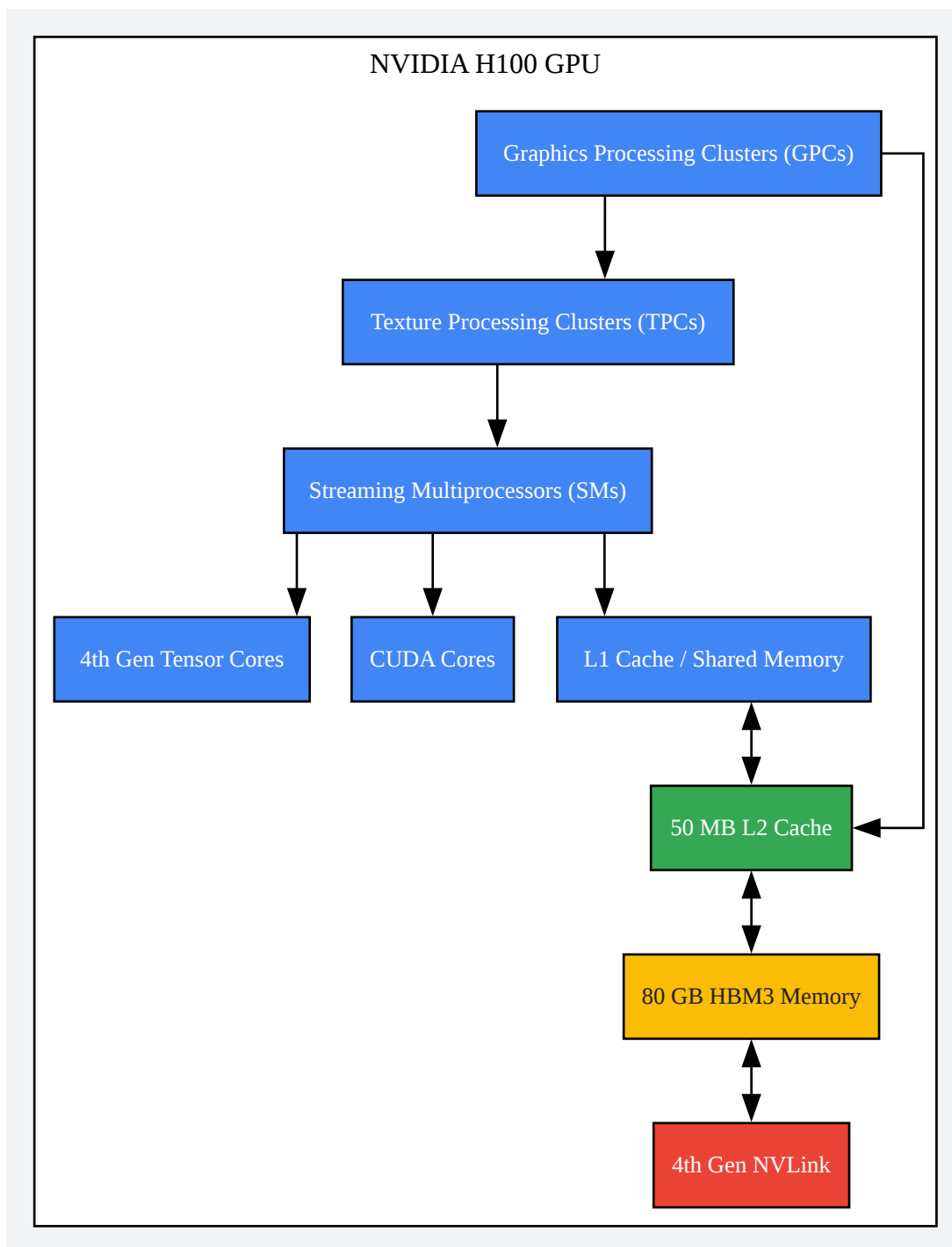
At the heart of the **H100** are the fourth-generation Tensor Cores, which accelerate a wide range of matrix operations crucial for AI and HPC workloads.[4] A key innovation is the Transformer Engine, which, combined with the new FP8 data format, dramatically speeds up the training and inference of transformer models, a cornerstone of modern AI.[4][5] The **H100** also boasts a significantly larger L2 cache and utilizes High Bandwidth Memory 3 (HBM3), providing a substantial boost in memory bandwidth, which is critical for handling the massive datasets common in scientific research.[2][6]

Architectural Specifications

The following tables summarize the key specifications of the **H100** GPU and compare it with its predecessor, the NVIDIA A100.

Metric	NVIDIA H100 (SXM5)[3]	NVIDIA A100 (SXM4 80GB) [7]
GPU Architecture	NVIDIA Hopper	NVIDIA Ampere
Manufacturing Process	TSMC 4N	TSMC 7N
Transistors	80 Billion	54.2 Billion
CUDA Cores	16,896	6,912
Tensor Cores	528 (4th Gen)	432 (3rd Gen)
GPU Boost Clock	1.78 GHz	1.41 GHz
GPU Memory	80 GB HBM3	80 GB HBM2e
Memory Bandwidth	3 TB/s	2 TB/s
L2 Cache	50 MB	40 MB
NVLink Interconnect	900 GB/s (4th Gen)	600 GB/s (3rd Gen)
TDP	700W	400W
Peak Performance	NVIDIA H100[7]	NVIDIA A100 (80GB)[7]
FP64	30 TFLOPS	9.7 TFLOPS
FP64 Tensor Core	60 TFLOPS	19.5 TFLOPS
FP32	60 TFLOPS	19.5 TFLOPS
TF32 Tensor Core	500 TFLOPS	156 TFLOPS
FP16 Tensor Core	1,000 TFLOPS	312 TFLOPS
INT8 Tensor Core	2,000 TOPS	624 TOPS

Below is a diagram illustrating the high-level architecture of the **H100** Tensor Core GPU.



[Click to download full resolution via product page](#)

High-level architecture of the NVIDIA **H100** GPU.

Molecular Dynamics Simulations

The enhanced computational power and memory bandwidth of the **H100** make it exceptionally well-suited for molecular dynamics (MD) simulations, a critical tool in drug discovery for

studying the physical movements of atoms and molecules.

Performance Benchmarks

The following table presents a comparison of **H100** performance against other NVIDIA GPUs for common MD simulation benchmarks using AMBER 22.[8][9] Performance is measured in nanoseconds per day (ns/day), where higher is better.

Benchmark System	Atom Count	Ensemble/Time step	H100 (ns/day)	A100 (ns/day)
JAC (DHFR)	23,558	NVE / 4fs	1479.32	1199.22
JAC (DHFR)	23,558	NPT / 4fs	1424.90	1194.50
Factor IX	90,906	NVE / 2fs	389.18	271.36
Factor IX	90,906	NPT / 2fs	357.88	248.65
Cellulose	408,609	NVE / 2fs	119.27	-
Cellulose	408,609	NPT / 2fs	108.91	-
STMV	1,067,095	NPT / 2fs	70.15	-

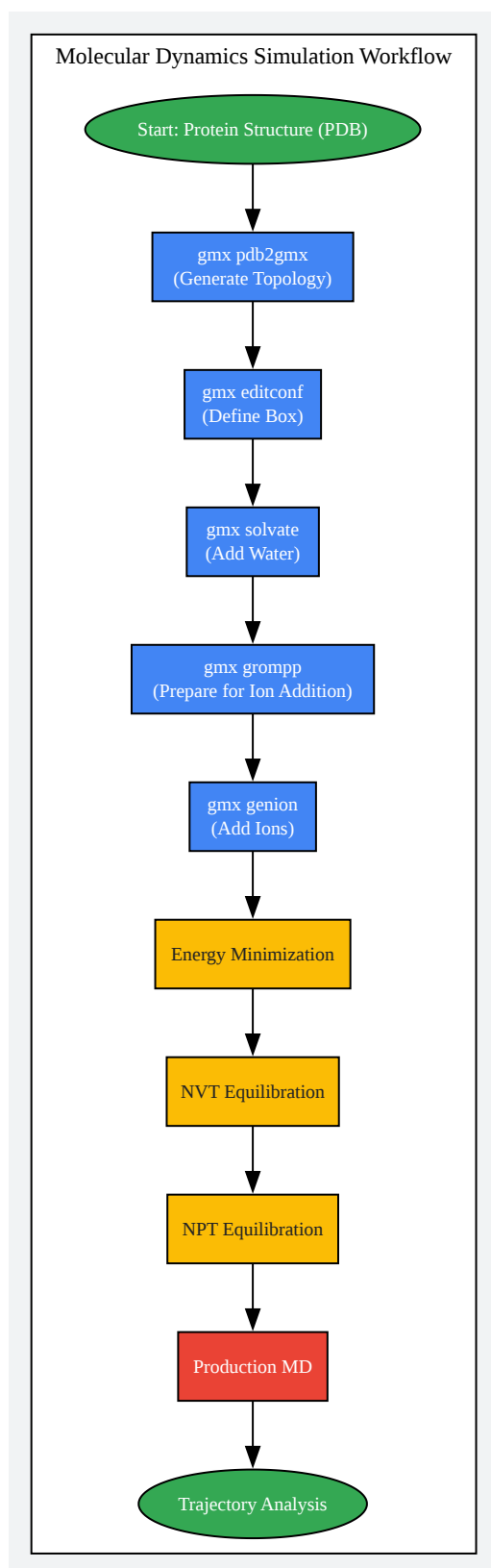
GROMACS benchmarks also demonstrate the **H100**'s superior performance, particularly for large systems.[10][11]

Benchmark System	Atom Count	H100 (ns/day)
System 1	~20,000	354.36
System 2 (RNA)	~32,000	1032.85
System 3 (Membrane Protein)	~80,000	400.43
System 4 (Protein in Water)	~170,000	204.96
System 5 (Membrane Channel)	~616,000	63.49
System 6 (Virus Protein)	~1,067,000	37.45

Experimental Protocol: Lysozyme in Water (GROMACS)

This protocol outlines the steps for a standard MD simulation of lysozyme in a water box using GROMACS.^{[12][13]}

- System Preparation:
 - Obtain the protein structure (e.g., PDB ID: 1AKI).
 - Generate a GROMACS topology using `gmx pdb2gmx`, selecting a force field (e.g., OPLS-AA/L all-atom force field) and water model (e.g., TIP3P).^[14]
 - Create a simulation box using `gmx editconf`, defining the box dimensions and centering the protein.
 - Fill the box with solvent (water) using `gmx solvate`.
 - Add ions to neutralize the system using `gmx grompp` and `gmx genion`.
- Energy Minimization:
 - Perform energy minimization to relax the system and remove steric clashes using `gmx grompp` and `gmx mdrun` with an appropriate `.mdp` file for minimization.
- Equilibration:
 - Conduct NVT (constant number of particles, volume, and temperature) equilibration to stabilize the temperature of the system.
 - Perform NPT (constant number of particles, pressure, and temperature) equilibration to stabilize the pressure and density.
- Production MD:
 - Run the production simulation for the desired length of time using `gmx mdrun`.



[Click to download full resolution via product page](#)

A typical workflow for a molecular dynamics simulation.

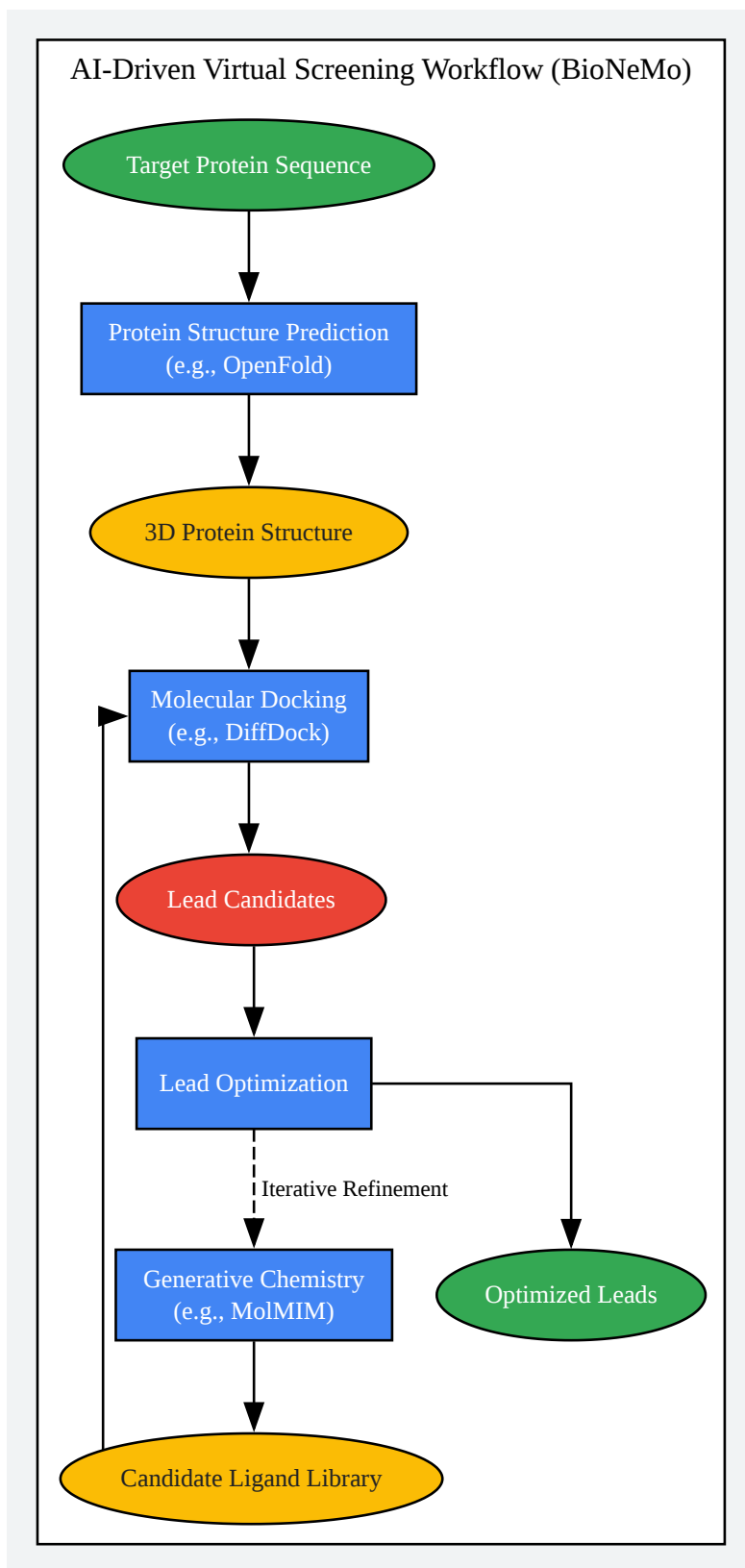
AI-Driven Drug Discovery

The **H100**'s prowess in AI computation is revolutionizing drug discovery by enabling large-scale virtual screening and generative chemistry. NVIDIA's BioNeMo platform provides a suite of pre-trained models and workflows optimized for the **H100**.[\[15\]](#)[\[16\]](#)

Experimental Protocol: Virtual Screening with BioNeMo

This protocol describes a generative virtual screening workflow using BioNeMo NIMs (NVIDIA Inference Microservices).[\[17\]](#)[\[18\]](#)

- Target Preparation:
 - Define the target protein sequence.
 - Use a protein structure prediction model like OpenFold to generate the 3D structure of the target protein.[\[17\]](#)
- Ligand Generation:
 - Employ a generative chemistry model, such as MolMIM, to generate a library of novel small molecules with desired properties.[\[18\]](#)
- Molecular Docking:
 - Utilize a docking tool like DiffDock to predict the binding affinity and pose of the generated ligands to the target protein.[\[19\]](#)
- Lead Optimization:
 - Analyze the docking results to identify promising lead candidates.
 - Iteratively refine the lead compounds by feeding the results back into the generative model for further optimization.



[Click to download full resolution via product page](#)

An AI-driven virtual screening workflow using BioNeMo.

Protein Structure Prediction

Predicting the 3D structure of proteins from their amino acid sequence is a computationally intensive task where the **H100** excels. Open-source models like OpenFold, a faithful reproduction of AlphaFold2, can be trained and run with significantly reduced time on **H100** GPUs.[\[20\]](#)

Performance Benchmarks

The training time for protein structure prediction models has been dramatically reduced with the **H100**.

Model/Framework	Hardware	Training Time
AlphaFold2 (Original)	128 TPUs	~11 days [20]
OpenFold (Generic)	128 A100 GPUs	> 8 days [20]
OpenFold (Optimized)	1,056 H100 GPUs	12.4 hours [20]
ScaleFold	2,080 H100 GPUs	10 hours [21] [22]

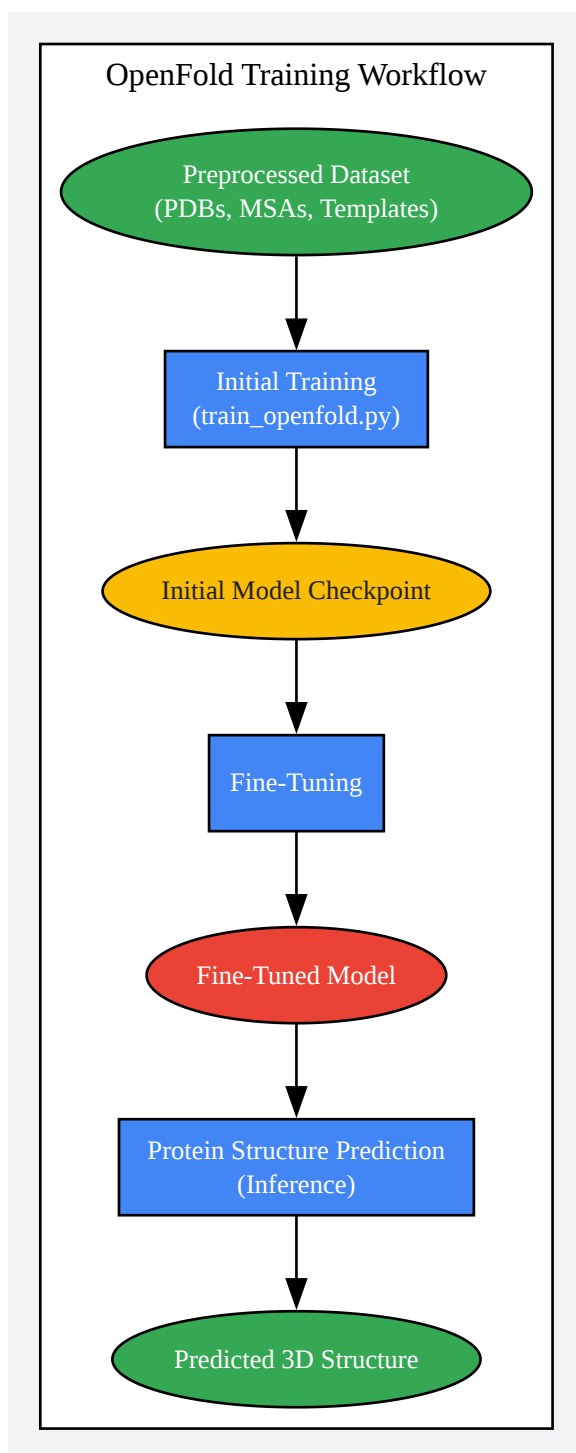
In the MLPerf HPC v3.0 benchmark, an optimized OpenFold implementation on **H100** GPUs completed a partial training task in just 7.51 minutes.[\[21\]](#)[\[22\]](#)

Experimental Protocol: Training OpenFold

This protocol outlines the general steps for training an OpenFold model.[\[23\]](#)

- Prerequisites:
 - Install OpenFold and its dependencies.
 - Prepare a preprocessed dataset of protein structures and their corresponding sequence alignments.
- Initial Training:

- Execute the `train_openfold.py` script with the paths to the training data, alignment directories, and template files.
- Specify training parameters such as the configuration preset, number of nodes and GPUs, and a random seed. For example, a training run might use a batch size of 128 for the initial 5,000 steps on 1,056 **H100** GPUs.[\[20\]](#)
- Fine-Tuning:
 - After the initial training phase, the model can be fine-tuned with different hyperparameters, such as an increased batch size, to further improve accuracy.[\[20\]](#)



[Click to download full resolution via product page](#)

A simplified workflow for training an OpenFold model.

Conclusion

The NVIDIA **H100** Tensor Core GPU, with its revolutionary Hopper architecture, provides researchers, scientists, and drug development professionals with a powerful tool to accelerate their work. From elucidating the dynamics of molecular interactions to designing novel therapeutics with generative AI and predicting the intricate structures of proteins, the **H100** is poised to drive the next wave of scientific breakthroughs. By understanding its architecture and leveraging optimized software and workflows, the research community can unlock new possibilities in the quest for novel medicines and a deeper understanding of biological systems.

Need Custom Synthesis?

BenchChem offers custom synthesis for rare earth carbides and specific isotopic labeling.

Email: info@benchchem.com or [Request Quote Online](#).

References

- [1. manuals.plus \[manuals.plus\]](#)
- [2. Hopper \(microarchitecture\) - Wikipedia \[en.wikipedia.org\]](#)
- [3. medium.com \[medium.com\]](#)
- [4. NVIDIA Hopper Architecture In-Depth | NVIDIA Technical Blog \[developer.nvidia.com\]](#)
- [5. lists.riscv.org \[lists.riscv.org\]](#)
- [6. ori.co \[ori.co\]](#)
- [7. H100 white paper - Jingchao's Website \[jingchaozhang.github.io\]](#)
- [8. \[AMBER\] H100 & RTX4090 Benchmarks from Ross Walker via AMBER on 2022-12-03 \(Amber Archive Dec 2022\) \[archive.ambermd.org\]](#)
- [9. Re: \[AMBER\] Amber 22 Nvidia GPU Benchmarks from Ross Walker via AMBER on 2023-03-21 \(Amber Archive Mar 2023\) \[archive.ambermd.org\]](#)
- [10. hpc.fau.de \[hpc.fau.de\]](#)
- [11. GROMACS 2024.1 on brand-new GPGPUs - NHR@FAU \[hpc.fau.de\]](#)
- [12. GROMACS Tutorials \[mdtutorials.com\]](#)
- [13. Lysozyme in Water \[mdtutorials.com\]](#)
- [14. m.youtube.com \[m.youtube.com\]](#)

- 15. Lysozyme in Water with aiida-gromacs - aiida-gromacs 2.1.1 documentation [aiida-gromacs.readthedocs.io]
- 16. BioNeMo for Biopharma | Drug Discovery with Generative AI | NVIDIA [nvidia.com]
- 17. GitHub - NVIDIA-BioNeMo-blueprints/generative-virtual-screening: NVIDIA BioNeMo blueprint for generative AI-based virtual screening [github.com]
- 18. youtube.com [youtube.com]
- 19. m.youtube.com [m.youtube.com]
- 20. Optimizing OpenFold Training for Drug Discovery | NVIDIA Technical Blog [developer.nvidia.com]
- 21. ScaleFold: Reducing AlphaFold Initial Training Time to 10 Hours [arxiv.org]
- 22. arxiv.org [arxiv.org]
- 23. Training OpenFold - OpenFold documentation [openfold.readthedocs.io]
- To cite this document: BenchChem. [H100 Tensor Core Architecture: A New Paradigm in Computation]. BenchChem, [2026]. [Online PDF]. Available at: [https://www.benchchem.com/product/b15585327/docs#h100-tensor-core-architecture-a-new-paradigm-in-computation]

Disclaimer & Data Validity:

The information provided in this document is for Research Use Only (RUO) and is strictly not intended for diagnostic or therapeutic procedures. While BenchChem strives to provide accurate protocols, we make no warranties, express or implied, regarding the fitness of this product for every specific experimental setup.

Technical Support: The protocols provided are for reference purposes. Unsure if this reagent suits your experiment?

Need Industrial/Bulk Grade? [Request Custom Synthesis Quote](#)

BenchChem

Our mission is to be the trusted global source of essential and advanced chemicals, empowering scientists and researchers to drive progress in science and industry.

Contact

Address: 3281 E Guasti Rd

Ontario, CA 91761, United States

Phone: (601) 213-4426

Email: info@benchchem.com

[Contact our Ph.D. Support Team for a compatibility check](#)