

Key features of NVIDIA H100 for academic research

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

Compound of Interest

Compound Name: H100

Cat. No.: B15585327

[Get Quote](#)

An In-Depth Technical Guide to the NVIDIA H100 for Academic Research

For researchers, scientists, and professionals in drug development, the NVIDIA H100 Tensor Core GPU, based on the Hopper architecture, represents a significant leap in computational power. This guide delves into the core technical features of the H100, providing a comprehensive overview of its capabilities for demanding academic research workloads.

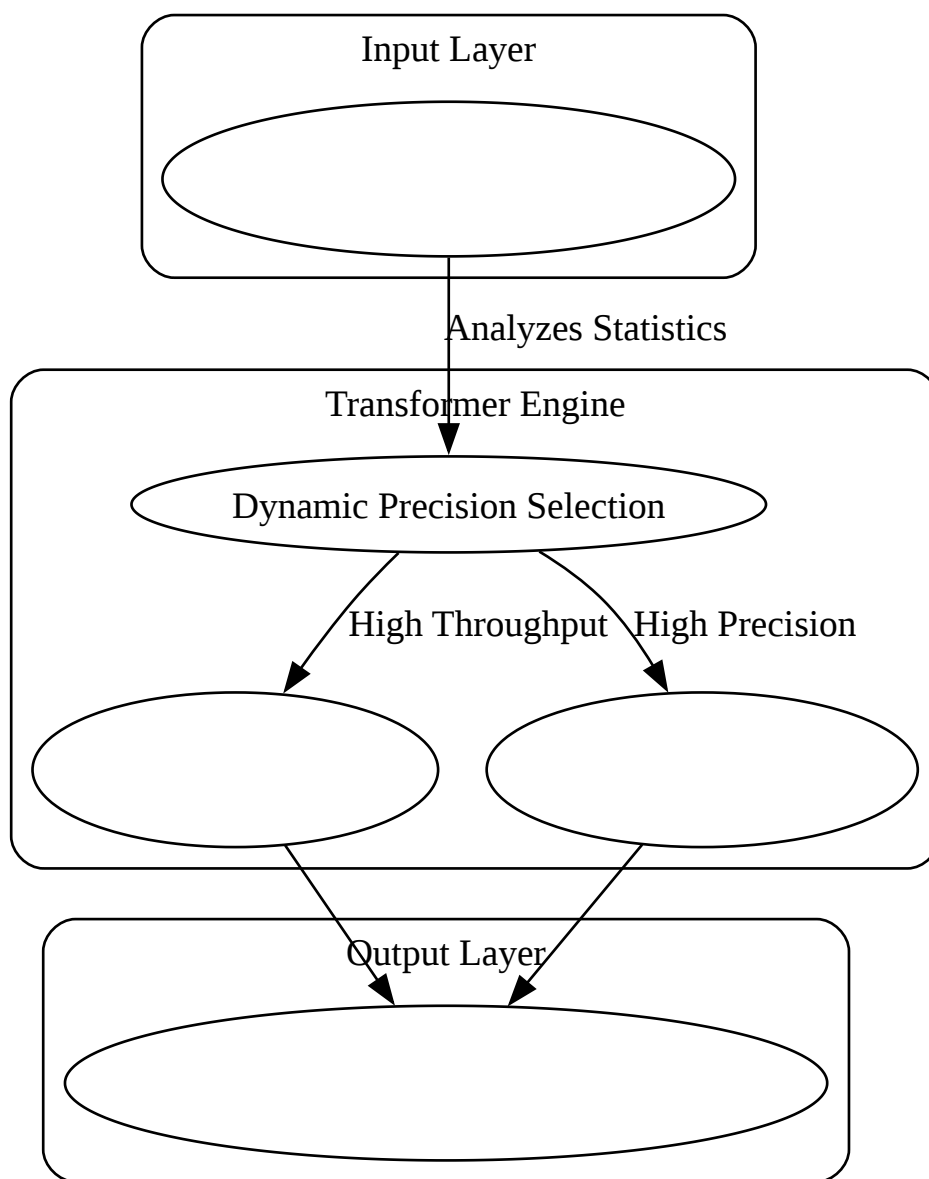
Architectural Innovations of the Hopper GPU

The NVIDIA H100 is built upon the Hopper architecture, which introduces several key advancements over its predecessors. Manufactured using a custom TSMC 4N process, it packs 80 billion transistors, leading to substantial improvements in performance and efficiency. [\[1\]](#)[\[2\]](#)

Fourth-Generation Tensor Cores and the Transformer Engine

At the heart of the H100's AI prowess are its fourth-generation Tensor Cores. These cores, combined with the innovative Transformer Engine, deliver unprecedented acceleration for AI workloads, particularly for the large language models (LLMs) and transformer networks that are

increasingly used in scientific research.[3] The Transformer Engine intelligently manages and dynamically selects between FP8 and 16-bit calculations, which can lead to up to 9 times faster AI training and 30 times faster AI inference on large language models compared to the A100.[4]

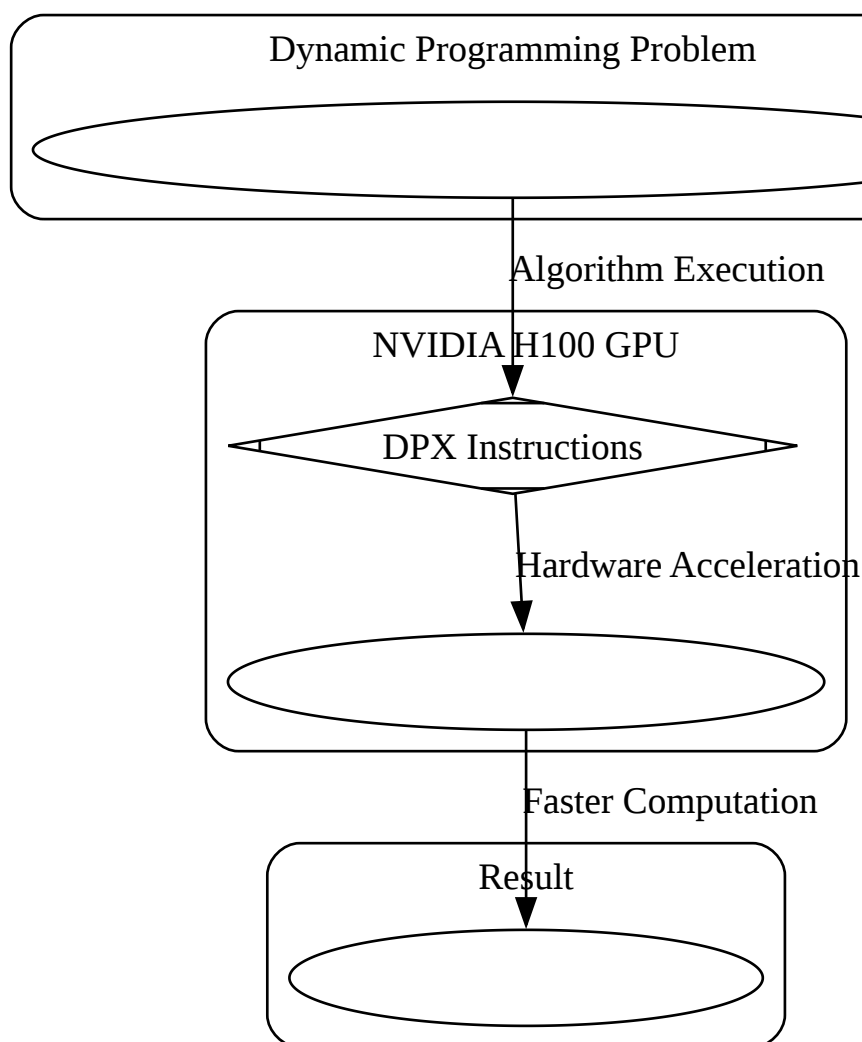


[Click to download full resolution via product page](#)

DPX Instructions for Dynamic Programming

The **H100** introduces a new instruction set, DPX, designed to accelerate dynamic programming algorithms.[5] This is particularly beneficial for bioinformatics research, such as DNA sequence alignment using algorithms like Smith-Waterman and Needleman-Wunsch.[1][5] With DPX

instructions, these algorithms can see speedups of up to 7 times compared to the NVIDIA A100 GPU.[5] In a multi-GPU setup, the acceleration can be even more significant, with a node of four **H100** GPUs speeding up the Smith-Waterman algorithm by 35 times.[1]



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

Quantitative Specifications

The following tables summarize the key quantitative specifications of the NVIDIA **H100**, with comparisons to its predecessor, the NVIDIA A100, where relevant.

Table 1: Core Architectural Specifications

Feature	NVIDIA H100 (SXM5)	NVIDIA H100 (PCIe)	NVIDIA A100 (SXM4)
GPU Architecture	Hopper	Hopper	Ampere
Transistors	80 Billion	80 Billion	54 Billion
CUDA Cores	16,896	14,592	6,912
Tensor Cores	528 (4th Gen)	456 (4th Gen)	432 (3rd Gen)
L2 Cache	50 MB	50 MB	40 MB
Memory	80 GB HBM3	80 GB HBM2e	80 GB HBM2e
Memory Bandwidth	3.35 TB/s	2 TB/s	2 TB/s
NVLink	4th Gen (900 GB/s)	4th Gen (600 GB/s via bridge)	3rd Gen (600 GB/s)
TDP	Up to 700W	350W	400W

Sources:[\[1\]](#)[\[2\]](#)[\[6\]](#)[\[7\]](#)[\[8\]](#)[\[9\]](#)

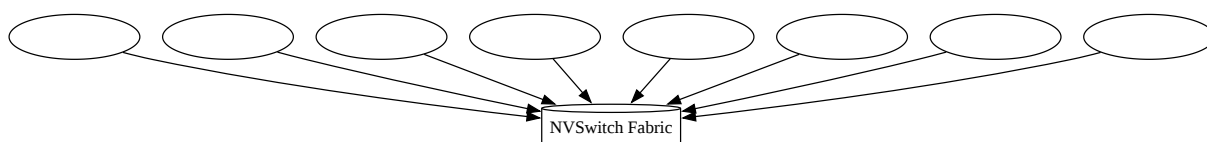
Table 2: Peak Performance (TFLOPS)

Precision	H100 SXM5	H100 PCIe	A100 SXM4
FP64	34	26	9.7
FP64 Tensor Core	67	51	19.5
FP32	67	51	19.5
TF32 Tensor Core	989	756	312
BFLOAT16 Tensor Core	1,979	1,513	624
FP16 Tensor Core	1,979	1,513	624
FP8 Tensor Core	3,958	3,026	N/A
INT8 Tensor Core	3,958	3,026	1,248

*With Sparsity. Sources:[4][10][11]

Multi-GPU Scalability with NVLink and NVSwitch

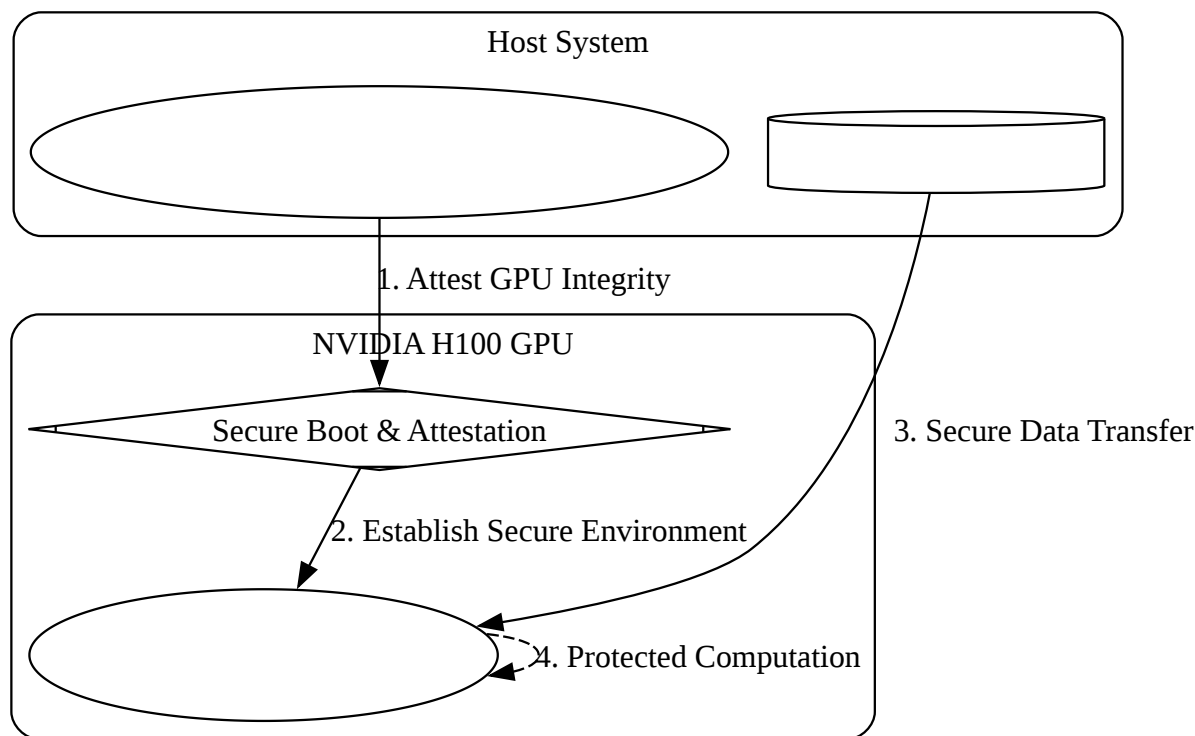
For large-scale research problems, the ability to scale across multiple GPUs is critical. The **H100** features fourth-generation NVLink, providing up to 900 GB/s of GPU-to-GPU interconnect bandwidth.[12] This is further enhanced by the third-generation NVSwitch technology, which allows for the creation of large, multi-node GPU clusters.[13] An NVLink Switch System can connect up to 256 **H100** GPUs, enabling massive parallel processing for the most demanding computational tasks.[4][13]



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

Confidential Computing for Secure Research

A groundbreaking feature of the **H100** is its support for confidential computing, making it the world's first accelerator with this capability.[4][14] This is crucial for research involving sensitive data, such as patient records in drug discovery or proprietary genomic data. Confidential computing on the **H100** creates a hardware-based trusted execution environment (TEE), isolating and protecting data and code in use, even from the host operating system or hypervisor.[14][15] The process involves a secure boot, hardware-level memory encryption, and an attestation mechanism to verify the integrity of the environment before processing sensitive data.[3][15]



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

Experimental Protocols for Key Research Areas

The following sections provide detailed methodologies for benchmark experiments in key academic research domains, based on published studies and best practices.

Molecular Dynamics Simulations

Objective: To benchmark the performance of molecular dynamics simulations using software packages like GROMACS or AMBER on the NVIDIA **H100**.

Methodology:

- System Preparation:
 - Select a standard benchmark system, such as the STMV (Satellite Tobacco Mosaic Virus) with approximately 1 million atoms, or a smaller system like DHFR (Dihydrofolate

reductase) in explicit solvent.

- Prepare the system using standard molecular dynamics setup procedures, including solvation, ionization, and energy minimization.
- Simulation Parameters (for GROMACS):
 - Integrator:md (leap-frog)
 - Time Step: 2 fs
 - Coulomb Type: PME (Particle Mesh Ewald)
 - Cut-off Scheme: Verlet
 - Constraints:h-bonds
 - Temperature Coupling: V-rescale
 - Pressure Coupling: Parrinello-Rahman
- Execution:
 - Run the simulation on a single **H100** GPU and record the simulation performance in nanoseconds per day (ns/day).
 - For multi-GPU benchmarks, utilize a system with multiple **H100s** connected via NVLink and run the simulation in parallel, ensuring the workload is appropriately distributed.
 - For enhanced throughput with multiple smaller simulations, NVIDIA Multi-Process Service (MPS) can be employed to allow multiple MPI ranks to run concurrently on a single GPU.
- Data Analysis:
 - Compare the ns/day performance of the **H100** against previous generation GPUs like the A100.
 - Analyze the scalability of performance with an increasing number of GPUs.

Sources:[13][16][17]

Genomics and Bioinformatics: Smith-Waterman Algorithm

Objective: To evaluate the performance of the Smith-Waterman algorithm for protein database search using the CUDASW++ software on the NVIDIA **H100**.

Methodology:

- Software and Database:
 - Use the CUDASW++ 4.0 software, which is optimized for modern GPU architectures and utilizes the **H100**'s DPX instructions.[14]
 - Benchmark against standard protein databases such as Swiss-Prot, UniRef50, or TrEMBL.[14]
- Execution Parameters:
 - Run the database search with a set of query protein sequences of varying lengths.
 - Utilize a standard substitution matrix, such as BLOSUM62.
 - Execute the search on a single **H100** GPU.
- Performance Measurement:
 - The primary performance metric is Giga Cell Updates Per Second (GCUPS), which measures the rate at which the dynamic programming matrix cells are computed.
 - Record the total execution time for the database search.
- Comparative Analysis:
 - Compare the GCUPS and execution time of the **H100** with the A100 and other GPU architectures.

- Evaluate the performance scaling when using multiple **H100** GPUs in a single node.

Sources:[5][14][15]

Deep Learning for Drug Discovery

Objective: To benchmark the training performance of a deep learning model for a drug discovery task, such as predicting protein-ligand binding affinity, on the NVIDIA **H100**.

Methodology:

- Model and Dataset:
 - Utilize a graph neural network (GNN) architecture, which is well-suited for learning from molecular structures.
 - Train the model on a large-scale public dataset such as PDBbind or BindingDB, which contain protein-ligand complexes with experimentally determined binding affinities.
- Training Protocol:
 - Framework: Use a popular deep learning framework like PyTorch or TensorFlow.
 - Precision: Leverage the **H100**'s Transformer Engine to enable mixed-precision training with FP8 for accelerated performance.
 - Hyperparameters:
 - Batch Size: Maximize the batch size that fits into the **H100**'s 80 GB HBM3 memory.
 - Optimizer: Adam or AdamW.
 - Learning Rate: Use a learning rate scheduler, such as cosine annealing.
 - Distributed Training: For multi-GPU training, use a library like PyTorch DistributedDataParallel (DDP) to scale the training across multiple **H100** GPUs.
- Performance Evaluation:

- Measure the training throughput in terms of the number of molecules processed per second.
- Record the time to convergence to a target validation accuracy or loss.
- Benchmarking:
 - Compare the training throughput and time to convergence of the **H100** with the A100.
 - Analyze the impact of using FP8 precision on both training speed and final model accuracy.

Sources:[10]

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. H100 Transformer Engine Supercharges AI Training, Delivering Up to 6x Higher Performance Without Losing Accuracy | NVIDIA Blog \[blogs.nvidia.com\]](#)
- [2. Nvidia's Next GPU Shows That Transformers Are Transforming AI - IEEE Spectrum \[spectrum.ieee.org\]](#)
- [3. forbes.com \[forbes.com\]](#)
- [4. lifescisearch.com \[lifescisearch.com\]](#)
- [5. Improving the Mapping of Smith-Waterman Sequence Database Searches onto CUDA-Enabled GPUs - PMC \[pmc.ncbi.nlm.nih.gov\]](#)
- [6. lambda.ai \[lambda.ai\]](#)
- [7. How does the Transformer Engine improve performance in AI workloads? - Massed Compute \[massedcompute.com\]](#)
- [8. www2.eecs.berkeley.edu \[www2.eecs.berkeley.edu\]](#)
- [9. easychair.org \[easychair.org\]](#)

- [10. newsletter.semianalysis.com \[newsletter.semianalysis.com\]](https://newsletter.semianalysis.com)
- [11. \[2406.10362\] A Comparison of the Performance of the Molecular Dynamics Simulation Package GROMACS Implemented in the SYCL and CUDA Programming Models \[arxiv.org\]](https://arxiv.org/abs/2406.10362)
- [12. CUDASW++4.0: ultra-fast GPU-based Smith-Waterman protein sequence database search - PubMed \[pubmed.ncbi.nlm.nih.gov\]](https://pubmed.ncbi.nlm.nih.gov/)
- [13. researchgate.net \[researchgate.net\]](https://researchgate.net)
- [14. cug.org \[cug.org\]](https://cug.org)
- [15. biorxiv.org \[biorxiv.org\]](https://biorxiv.org)
- [16. High Throughput AI-Driven Drug Discovery Pipeline | NVIDIA Technical Blog \[developer.nvidia.com\]](https://developer.nvidia.com)
- [17. intuitionlabs.ai \[intuitionlabs.ai\]](https://intuitionlabs.ai)
- To cite this document: BenchChem. [Key features of NVIDIA H100 for academic research]. BenchChem, [2026]. [Online PDF]. Available at: [\[https://www.benchchem.com/product/b15585327/docs#key-features-of-nvidia-h100-for-academic-research\]](https://www.benchchem.com/product/b15585327/docs#key-features-of-nvidia-h100-for-academic-research)

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](#)