

A Guide to HPC Resource Estimation

1. Understanding CPU and GPU Resource Usage

CPU Core-Hours

- Definition: 1 CPU core-hour = One CPU core being used for an hour.
- Examples of Usage: **256 CPU core-hours can be:**
- *256 CPU cores used for 1 hours.*
- *128 CPU cores used for 2 hour.*

GPU Card-Hours

- Definition: 1 GPU card-hour = One GPU card being used/reserved for an hour.
- Examples of Usage: **16 GPU card-hours can be:**
- *16 GPUs used for 1 hours.*
- *4 GPUs used for 4 hours.*
- *Note: 16 or 14 CPU core-hours comes free with each GPU card-hour requested in the same GPU node on ASPIRE 2A and ASPIRE 2A+ respectively.*

2. Memory – Examples of Capacity Limitations

- Training on large language models on GPU → Limited by VRAM.
- Quantum simulations on CPU or GPU → Limited by system RAM or VRAM, depending on implementation.

3. Storage

- Several types of storages in HPC: Parallel file systems, NVMe
- Project storage unit is in GB (Gigabyte)
- IO-Bound workload: Use high-speed scratch storage, parallel I/O (HDF5, MPI-IO)
- Key considerations: Checkpointing and data management

1. **Critical Role of Hardware:** Accurate resource estimation hinges on understanding the hardware available.
2. **Review Available Systems:** Familiarise yourself with the specifications of [ASPIRE 2A](#) and [ASPIRE 2A+](#) prior to estimation, as HPC application performance is closely tied to hardware details.
3. **System Allocation:** Note that applicants do not choose between ASPIRE 2A and ASPIRE 2A+; NSCC assigns the most suitable system based on your project descriptions.
4. **GPU Card-Hour Basis:** GPU card-hour estimates in project applications should be based on Nvidia A100 or H100 specifications.
5. **Shared System Considerations:** In most cases, avoid resource estimations that assume exclusive access to all specified hardware (e.g., planning for 256 GPUs to be simultaneously available for a single AI workload). Please write to us separately if you have such requirements.

Overview of NSCC Singapore's Supercomputers

>3.5x
More Cores

5x
More Compact



2x
More GPUs

7x
More Compute
Power

**Comparisons to ASPIRE 1 supercomputer*

105,984 Cores	352 GPUs	476TB	25 PBytes	10 PBytes
CPU (AMD EPYC™ 7713) 800 Nodes	Accelerated Nodes GPU (NVIDIA A100) 82 Nodes	Total System Memory	Storage (Spinning + Nearline)	Scratch Disk

HPC Top 500 Ranking	#1 	#2 	#3 	#4 	#90 	#301 	#429 
System	El Capitan	Frontier	Aurora	Eagle	ASPIRE 2A+	ASPIRE 2A (GPU)	ASPIRE 2A (CPU)
Manufactured by	HPE Cray	HPE Cray	HPE Cray	Microsoft	Nvidia DGX	HPE Cray	HPE Cray
Total Pflops	1,742.00	1,353.00	1012.10	561.70	14.20	3.33	2.58

**Ranking as of Nov 2024*

Knowing Your Hardware (ASPIRE 2A)

Full list of HPC cluster node types and specification:

Server	CPU Model	Cores per socket	Socket per server	Total Physical cores per server	Available RAM (DDR4)	GPUs
Standard Compute Node (768 nodes)	Dual-CPU AMD EPYC 7713	64	2	128	512 GB	No GPU
GPU compute node (64 nodes)	Single-CPU AMD EPYC 7713	64	1	64	512 GB	4x Nvidia A100 40GB
GPU AI Node (12 nodes)	Single-CPU AMD EPYC 7713	64	1	64	512 GB	4x Nvidia A100 40GB (11TB nvme)
GPU AI Node (6 nodes)	Dual-CPU AMD EPYC 7713	64	2	128	1 TB	8x Nvidia A100 40GB (14TB nvme)
Large memory node (12 nodes)	Dual-CPU AMD EPYC 7713	64	2	128	2 TB	No GPU
Large memory node (4 nodes)	Dual-CPU AMD EPYC 7713	64	2	128	4 TB	No GPU

Storage

HOME+Project (GPFS) $\approx$ 25 PB

Scratch (Lustre) $\approx$ 10 PB

Interconnect (Slingshot 10)

1x100G Link (CPU node, Large Memory Node)

2x100G Link (GPU Node)

Free 16 CPU core-hours with every 1 GPU card-hour requested.

Knowing Your Hardware (ASPIRE 2A+)

ASPIRE 2A+ Components	Specification
Compute Nodes	NVIDIA DGX SuperPOD™ - 40 Nodes of DGX H100 (See below)
Interconnect	NVIDIA Quantum 2 based NDR InfiniBand
Storage	Scratch ~ 2.5 PB, Home ~ 27.5 PB

DGX H100	Specification
CPU	Dual Intel Xeon Platinum 8480C Processors Total Cores = 2 x 56 Cores = 112 Cores
System Memory	2 TB
GPU	8 x NVIDIA H100 GPUs
GPU Memory	640 GB (80 GB on each GPU card)
Storage	8 x 3.84 TB NVMe drives
Network	4 x OSFP ports for 8 x NVIDIA® ConnectX®-7 Single Port InfiniBand Cards 8 x 400 Gb/s InfiniBand
NVSwitch	4 x 4th generation NVLinks that provide 900 GB/s GPU-to-GPU bandwidth
Operating System	DGX OS (Ubuntu 22.04)
Performance	FP64 - 272 teraFLOPS, TF32 (Tensor Core) - 7.9 petaFLOPS, FP8 - 32 petaFLOPS

Free 14 CPU core-hours with every 1 GPU card-hour requested.

Maximum Resources Available

HPC System	Maximum Amount of Resources Per Year.
ASPIRE 2A – Cray CPU Nodes	Total CPU core-hours per year = 768 Nodes x 128 Cores x 24 Hours x 365 Days = 861 Million CPU Core-hours RIE 60% Allocation per year = 517 Million CPU Core-hours
ASPIRE 2A – Cray GPU Nodes + AI Nodes	Total GPU card-hours per year = 352 A100 GPUs x 24 Hours x 365 Days = 3.08 Million GPU Card-hours RIE 60% Allocation per year = 1.85 Million A100 GPU Card-hours
ASPIRE 2A+ – DGX SuperPOD	Total GPU card-hours per year = 320 H100 GPUs x 24 Hours x 365 Days = 2.80 Million GPU Card-hours RIE 60% Allocation per year = 1.68 Million H100 GPU Card-hours



Do not enter numbers in your project application that exceed the amounts shown above.

Important Note: Resource allocation is competitive and subject to availability. Approved projects will receive resources for their entire duration. Additionally, resource usage is chargeable based on the category you are applying for.

Knowing Your Hardware (Storage)

1. In call for projects, the storage you are applying is a storage space **shared by all users** in the project, under the mount point of
`/home/project/<project-id>`
2. Regardless of project allocation, each user receives a complimentary 50GB in their \$HOME directory or
`/home/users/<your_org>/<userid>`
3. Scratch storage is available at no cost, with a limit of 100TB per user, BUT it is subject to a 30-day purge policy. Scratch storage is intended for temporary staging purposes.
`/home/users/<your_org>/<userid>/scratch`

1. Eligibility & Resource Considerations

- ASPIRE 2A or ASPIRE 2A+ assignment is done by NSCC based on project requirement.
- For small/modest workloads, consider departmental or institutional resources.
- Include a **10% resource buffer** for debugging, restarts, and unforeseen issues.

2. CPU vs GPU Guidance

- **AI workloads:** GPUs almost a must; pre/post-processing often CPU-only.
- **HPC applications:** mostly CPU; GPU offloading requires explicit porting. Check latest HPC application versions for optimisations and GPU support.

3. Estimation & Benchmarking

- Use past performance data carefully; outdated software or hardware may mislead.
- Estimate CPU/GPU hours and storage accurately per workload.
- Benchmark and develop a workflow to improve efficiency and avoid resource wastage.

1. Training & Workshops

- Mandatory NSCC Introductory Workshops for 2A/2A+ users.
- HPC Clinic (quarterly) and Planning Workshop (post-project calls).
- Training enables faster time-to-solution and reduces resource wastage.

2. Resource Request Guidelines

- Avoid unjustified or imbalanced requests (e.g., GPU for non-GPU workloads or excessive storage/compute ratio).
- Align compute and storage requests with actual application needs.

3. Continuity & Handover

- Plan for human resource continuity to prevent disruptions.
- Ensure **proper handover** if key personnel leaves.

4. System Awareness

- Understand HPC hardware/software stack; system architecture affects performance.
- Attend training or consult documentation for guidance.

Nvidia Benchmark (AI)– GPU

- <https://developer.nvidia.com/deep-learning-performance-training-inference/training>
- <https://developer.nvidia.com/deep-learning-performance-training-inference/ai-inference>
- <https://developer.nvidia.com/deep-learning-performance-training-inference/conversational-ai>
- https://docs.nvidia.com/nemo-framework/user-guide/24.09/performance/performance_summary.html

AI Training or Fine-tuning Resource Guideline – GPU

- [DeepSeek-V3 Technical Report](#)
- https://docs.api.nvidia.com/nim/reference/meta-llama-3_1-405b
- [Understanding the Performance and Estimating the Cost of LLM Fine-Tuning](#)
- https://en.wikipedia.org/wiki/Neural_scaling_law

Nvidia Benchmark (HPC Applications) – GPU

- <https://developer.nvidia.com/hpc-application-performance>

GROMACS, AMBER, NAMD, LAMMPS (HPC Application) – CPU and GPU

- <https://www.hecbiosim.ac.uk/access-hpc/our-benchmark-results/archer2-benchmarks>
- <https://www.hecbiosim.ac.uk/access-hpc/our-benchmark-results/bede-benchmarks>

NSCC User Guides (PDF documents)

[ASPIRE 2A - https://help.nsc.sg/aspire2a/user-guide/](https://help.nsc.sg/aspire2a/user-guide/)

[ASPIRE 2A+ - https://help.nsc.sg/aspire2aplus/user-guide/](https://help.nsc.sg/aspire2aplus/user-guide/)

NSCC User Guides (Markdown - Will be updated progressively)

<https://nscsg.github.io/>



Join our Slack Channel by scanning the QR code.

Technical questions? Simply send an email to help@nsc.sg



National
Supercomputing
Centre

Thank You