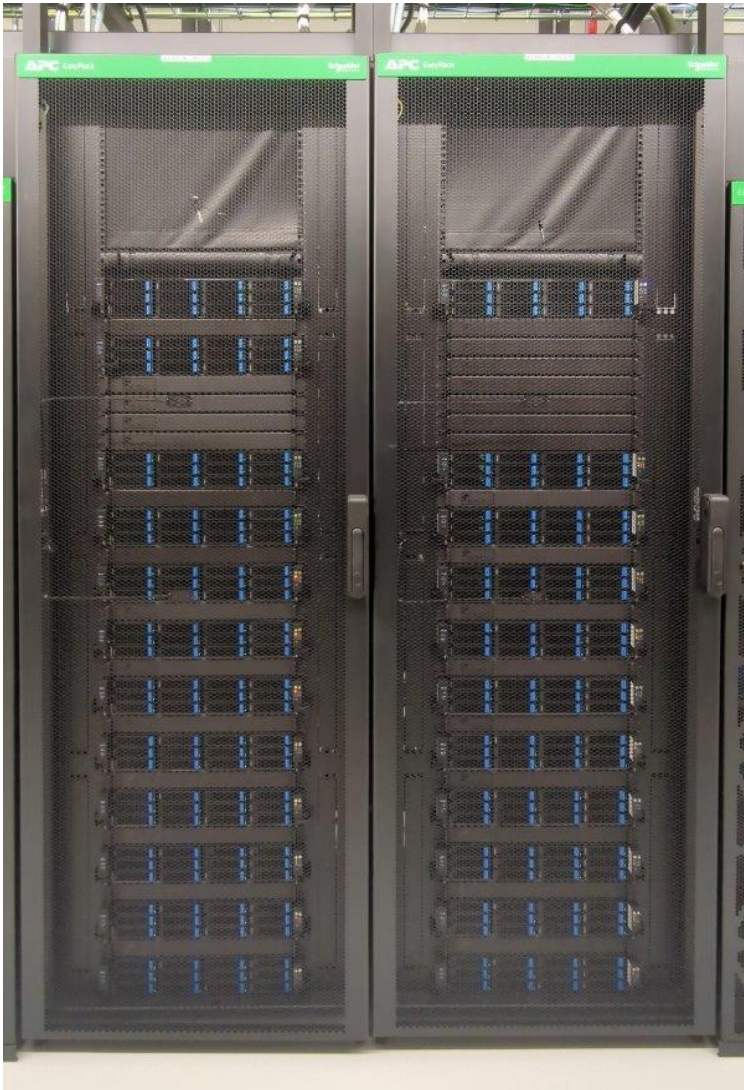


Welcome to Lyra

The new CÉCI-ULB GPU cluster

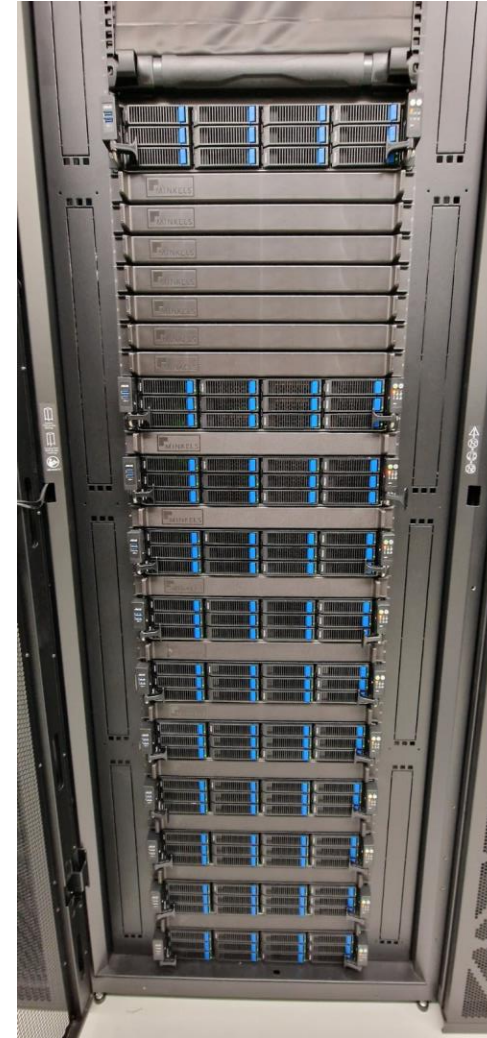
Nicolas Potvin
Ariel Lozano

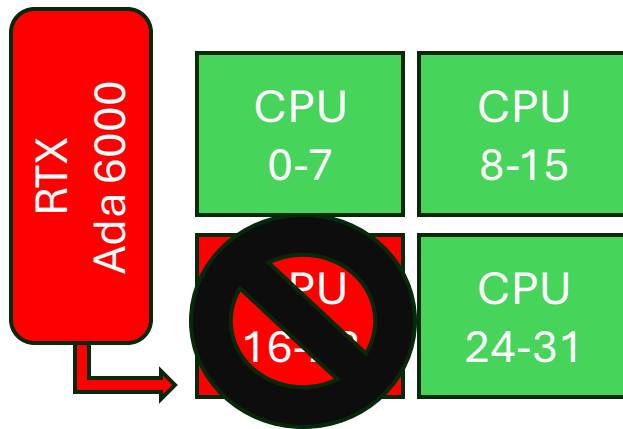




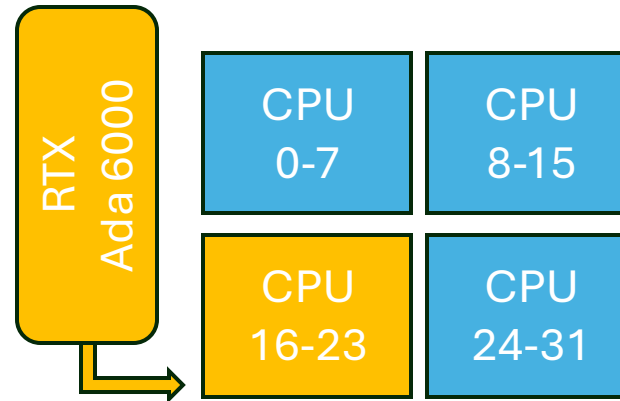
- ❖ 40 virtual Compute Nodes
 - 1x 32 cores 3.25GHz (x86_64-v4)
 - 128GB RAM
 - 1x Nvidia RTX 6000 ADA 48 GB
 - 1.5 TB LOCALSCRATCH
 - 50 GbE
- ❖ Totals:
 - 1280 cores
 - 5280 GB RAM
 - 40 GPUs
- ❖ Storage
 - 500 TB CephFS GLOBALSCRATCH
 - 30 TB CephFS home storage
 - CÉCI common storage available

- ❖ Favoured jobs
 - AI & Machine Learning
 - PyTorch, TensorFlow
- ❖ Welcomed jobs
 - GPU-ready applications
 - NAMD, GROMACS
 - Statistics, physics, *etc.*
- ❖ Also welcomed jobs
 - Serial (CPU) job arrays
 - SMP (OpenMP, *etc.*)
 - HTC

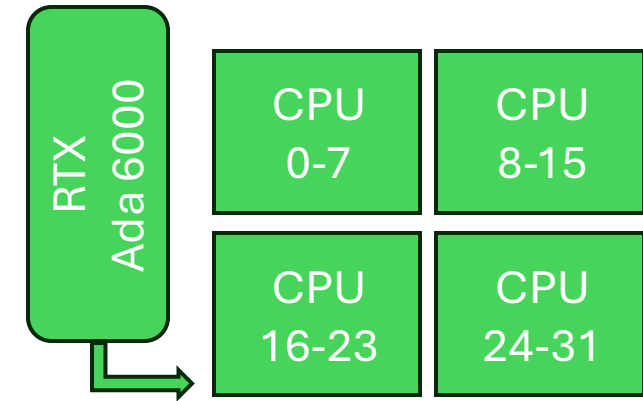




Ressources for CPU-only jobs



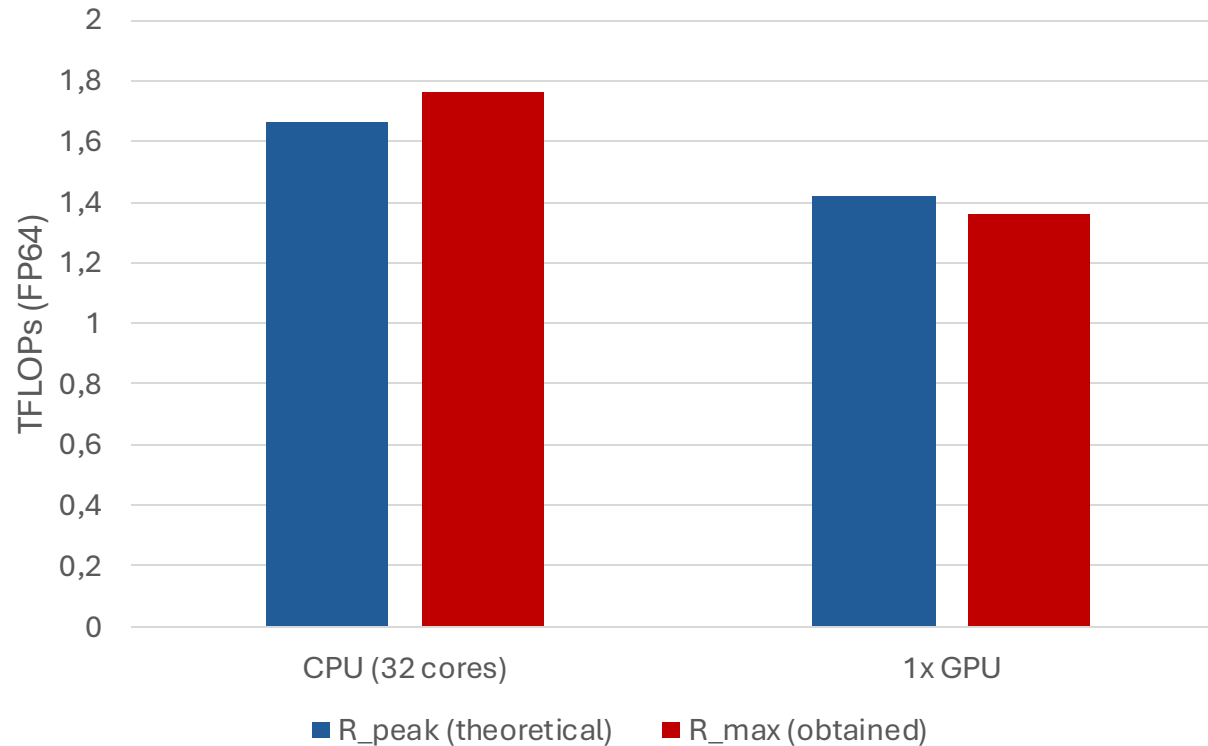
Ressources on a
compute node



Ressources for GPU-CPU jobs

- ❖ Login nodes
 - Equipped with a GPU
 - User RAM consumption limited to 8GB max
- ❖ Slurm scheduler configuration
 - Single *batch* partition
 - Max Walltime 5 days
 - 24 cores max without GPU (per node)
 - 8 cores reserved exclusively for GPU jobs
 - 256 cores and/or 8 GPUs max allocated per user
- ❖ Quotas
 - Home: 100GB / 100.000 files
 - GLOBALSCRATCH: 5TB / 300.000 files

HPL benchmarks



- ❖ Solves a dense system of linear equations in DP (FP64)
- ❖ CPU results
 - 106% score obtained!
 - Can be explained by turbo-boost
- ❖ GPU results
 - 96% score obtained
- ❖ The impact of virtualization is negligible (lot of fine-tuning work)
- ❖ Lyra provides 124.8 TFL0Ps of aggregated compute power

❖ Releases currently installed

- 2023a (default)
- 2023b, 2024a, 2025a
- More modules will be available soon

❖ Newer software

- Modules are preferred
- Apptainer (containers) when modules are unavailable
- Contact us for support (ceci-support@ulb.be)

❖ Apptainer container images

```

[17:48:35][alozano@lyra-l01]:~
>>$ srun --pty --cpus-per-task=4 --mem-per-cpu=8G bash -l

[17:48:38][alozano@ly-w109]:~
>> $ apptainer pull docker://pytorch/pytorch:2.7.0-cuda12.6-cudnn9-runtime
INFO:      Converting OCI blobs to SIF format
INFO:      Starting build...
INFO:      Fetching OCI image...
29.0MiB / 29.0MiB [=====] 100 % 0.0 b/s 0s
...
INFO:      Creating SIF file... [=====] 100 % 0s

[17:52:46][alozano@ly-w109]:~
>> $ ls
pytorch_2.7.0-cuda12.6-cudnn9-runtime.sif
<Ctrl+D>
[17:52:48][alozano@lyra-l01]:~
>> $ apptainer exec --nv pytorch_2.7.0-cuda12.6-cudnn9-runtime.sif python
Python 3.11.12 | packaged by conda-forge | (main, Apr 10 2025, 22:23:25) [GCC 13.3.0] on linux
Type "help", "copyright", "credits" or "license" for more information.
>>> import torch
>>> torch.cuda.get_device_name(0)
'NVIDIA RTX 6000 Ada Generation'
  
```

- ❖ Debug partition
 - 2 nodes
 - Walltime 30m
 - Builds, job submission validation, ...

- ❖ CernVM-FS distributed software
 - EESSI (European Environment for Scientific Software Installations)
 - Cern ATLAS

- ❖ The team making Lyra possible
 - Nicolas Potvin
 - Ariel Lozano
 - Michaël Waumans
 - Arthur Lesuisse
 - Jean-Luc Zirani
 - Raphaël Leplae

