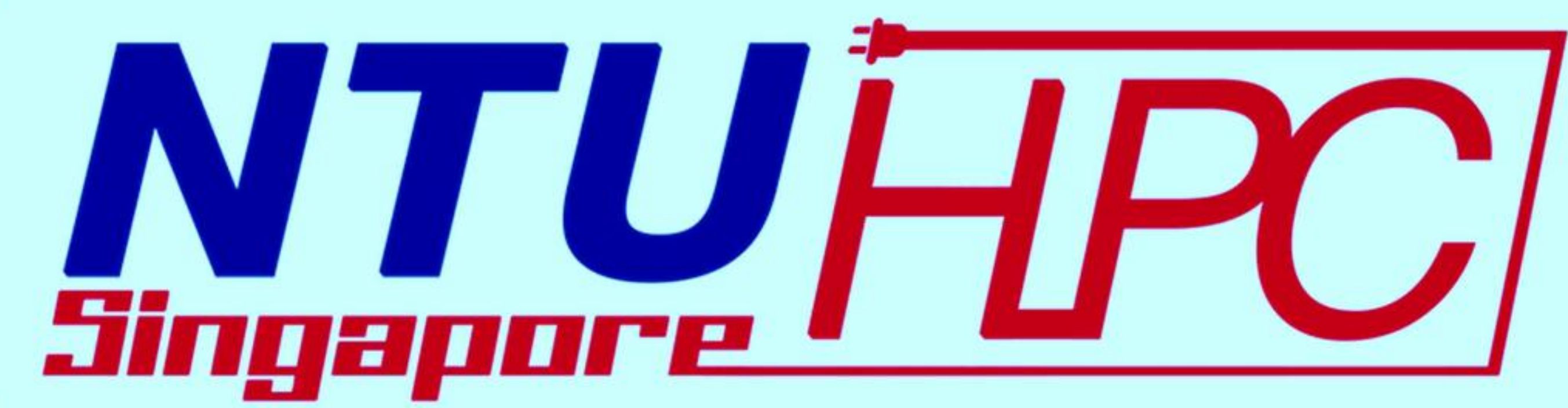


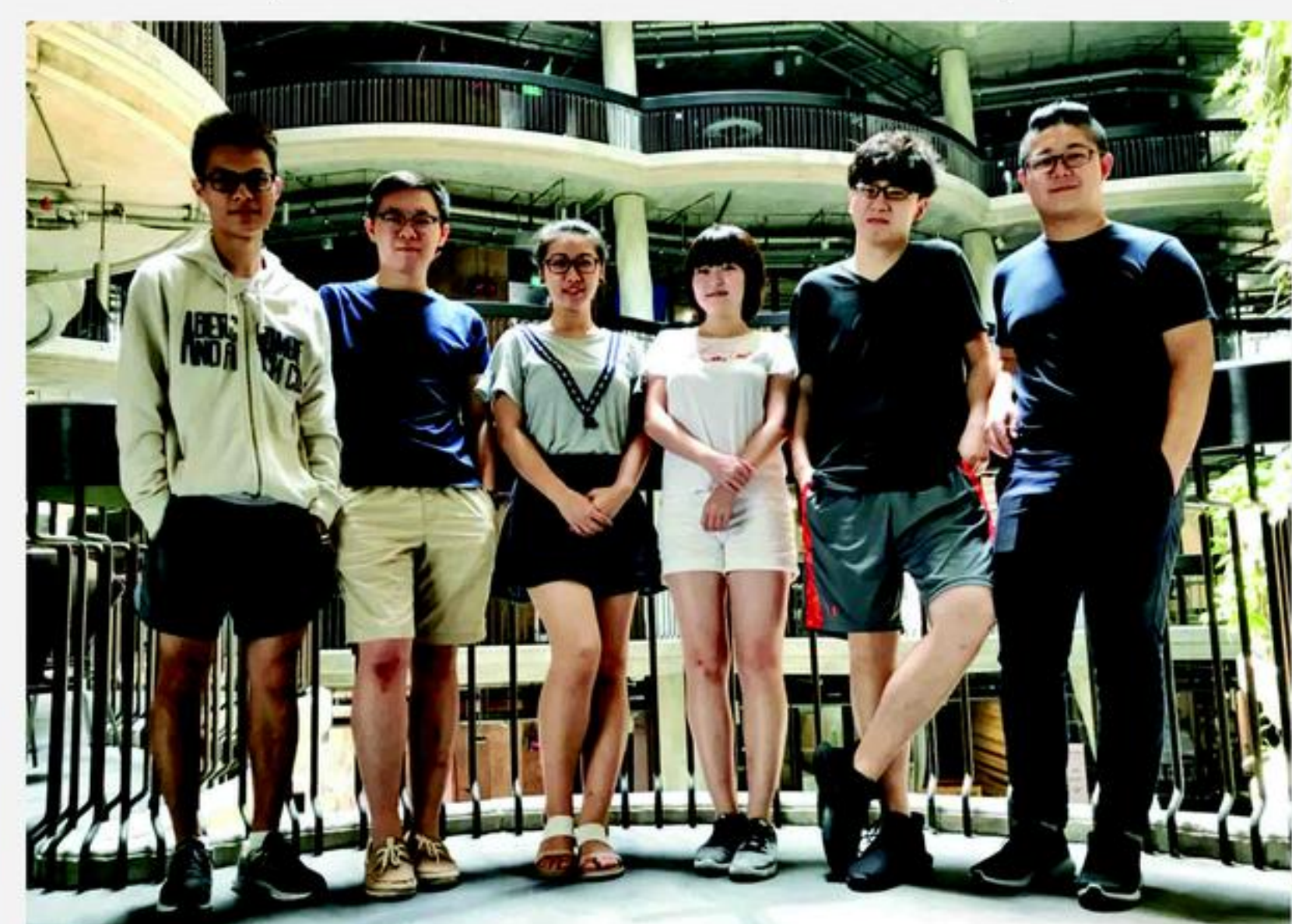
Team Supernova



Students: Siyuan Liu, Yiyang Shao, Hailin Chen, Ziji Shi, Meiru Hao, Shuqian Tang
Advisors: Prof. Bu-Sung Lee (NTU), Mr. Jeff Adie (NVIDIA), Mr. Jianxiong Yin (NVIDIA),
Dr. James Chen (NSCC), Mr. Paul Hiew (NSCC)

Team & School

Inaugurated in 1991, NTU has grown to become a comprehensive and research-intensive university. Our university is consistently ranked as one of the top universities in Asia. In NTU, the High Performance Computing Centre (HCCC) provides HPC resources to the entire NTU research community. What's more, NTU's School of Computer Science and Engineering offers a HPC specialization to computer science and computer engineering students.



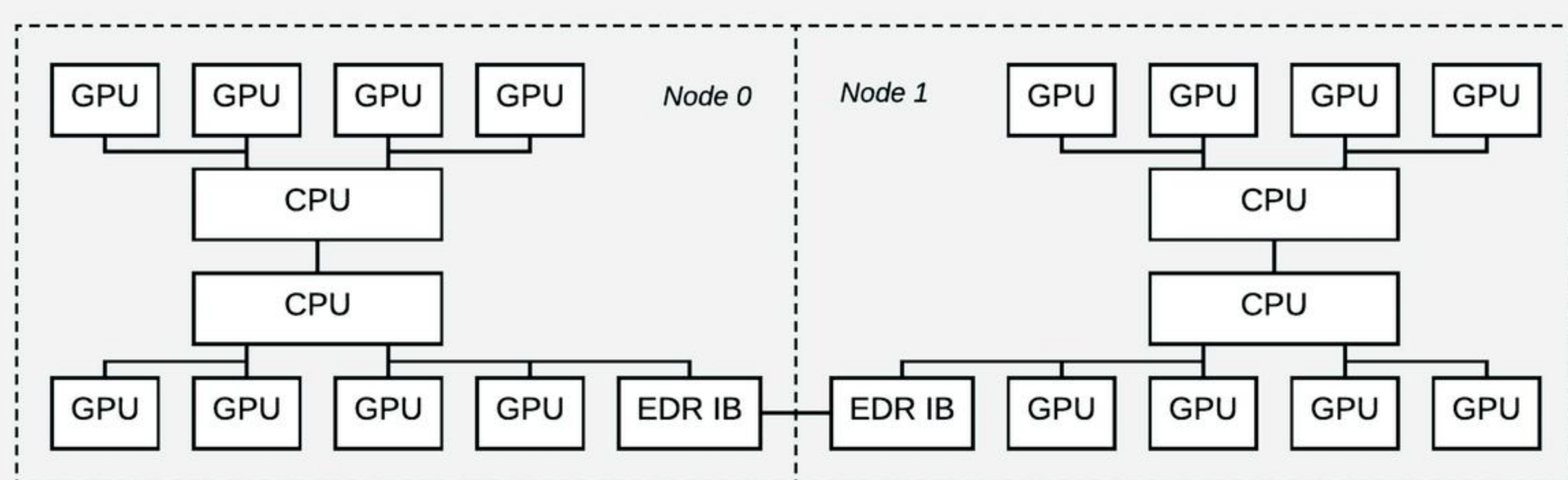
From left to right: Hailin, Ziji, Shuqian, Meiru, Yiyang, Siyuan.

Our team is a truly diverse team. The six members have a diverse background in computer science, including algorithm design, operating systems, artificial intelligence, and high performance computing. The male/female ratio in our team is also on par with national demographics. We together work hard to prepare for this competition.

Hardware & Software

We are sporting two Supermicro 4028GR-TR servers. Each of them has the following configuration:

CPU	Intel Xeon E5-2699 V4 (22 cores) * 2
GPU	NVIDIA Tesla V100 * 8
Memory	256GB 2133MHz DDR4
Interconnect	Mellanox InfiniBand 100Gb/s EDR
Storage	480GB SSD * 1 1TB SSD * 4 (RAID0 on master)



We designed a high-density GPU cluster because most of the applications (HPL, HPCG, Born) runs well on GPUs. For CPU applications, we can tackle them with 88 cores and the cloud component. Although theoretically the maximum power consumption of our cluster is way more than 3,000 Watts, for real world applications that we have tested (which are mostly memory bound), it runs well under the power limit. For compute-intensive applications, we can adjust the frequencies to lower the power consumption. Plus, we save the power of the InfiniBand switch by having just two nodes.

On the software side, we are using CentOS 7.3 with NVIDIA CUDA Toolkit 9.0, Intel Parallel Studio 2018, Mellanox HPC-X 1.9.7, and Alinea Forge 7. These software suites are well optimized for our hardware and enable us to gain the most performance out of our cluster. We also use utilities such as Lmod, EasyBuild, Ansible, and Grafana to better manage our cluster.

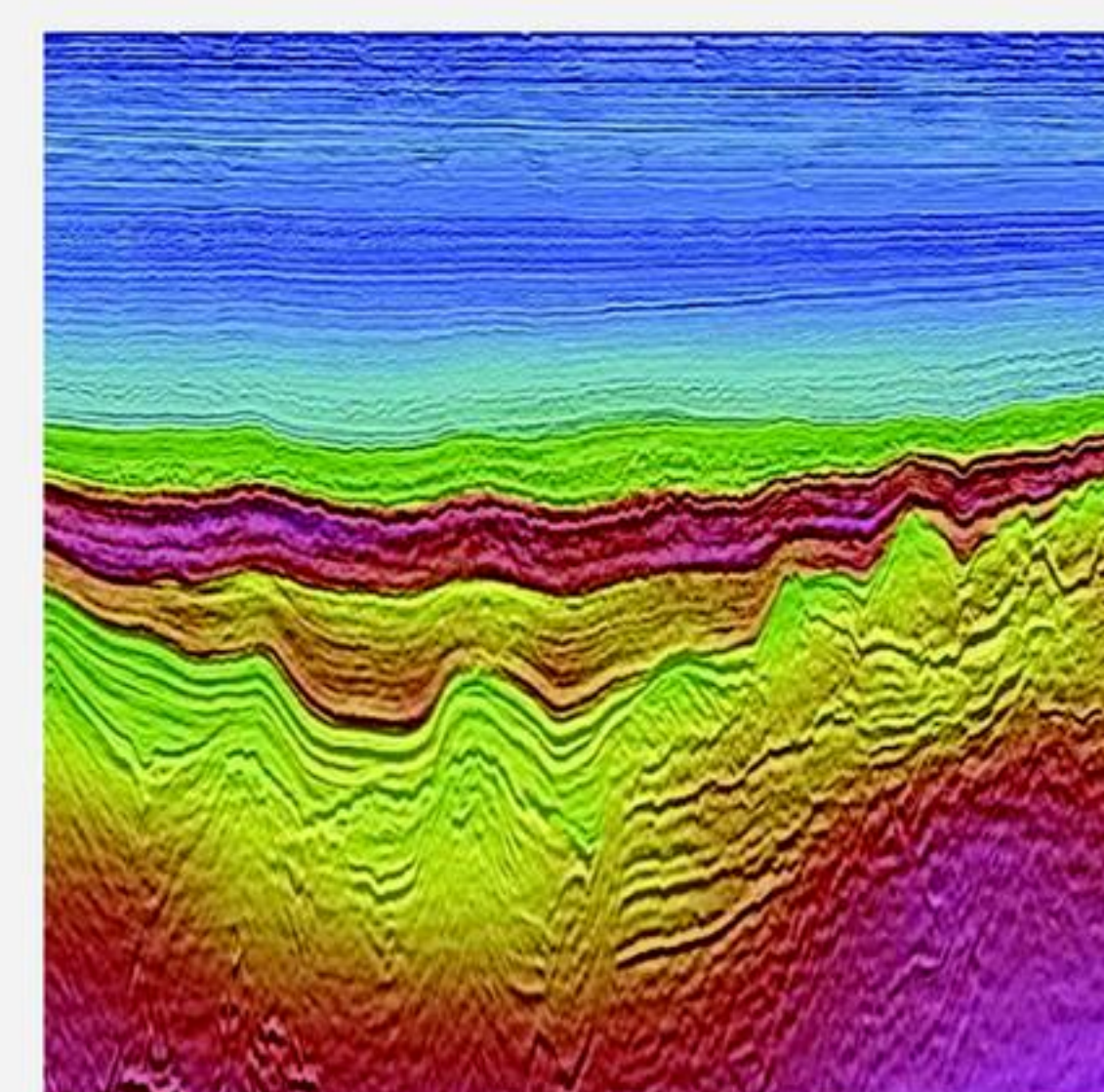
Preparations

1. We allocated each application to multiple members so that we can compliment each other and better brainstorm for optimization ideas.
2. We held weekly meeting to discuss applications and HPC technologies.
3. We consulted domain experts on seismic imaging, phylogenetics, molecular dynamics, and GPU to better understand the applications.
4. We attended some conferences and events, such as the Singapore Supercomputing Frontier conference, to accumulate experience in HPC.
5. We took the Parallel Computing course which includes MPI, OpenMP, and CUDA programming and performance analysis.
6. We maintain a small training cluster with multiple nodes and GPU support for practising before competition hardware arrives.

Applications

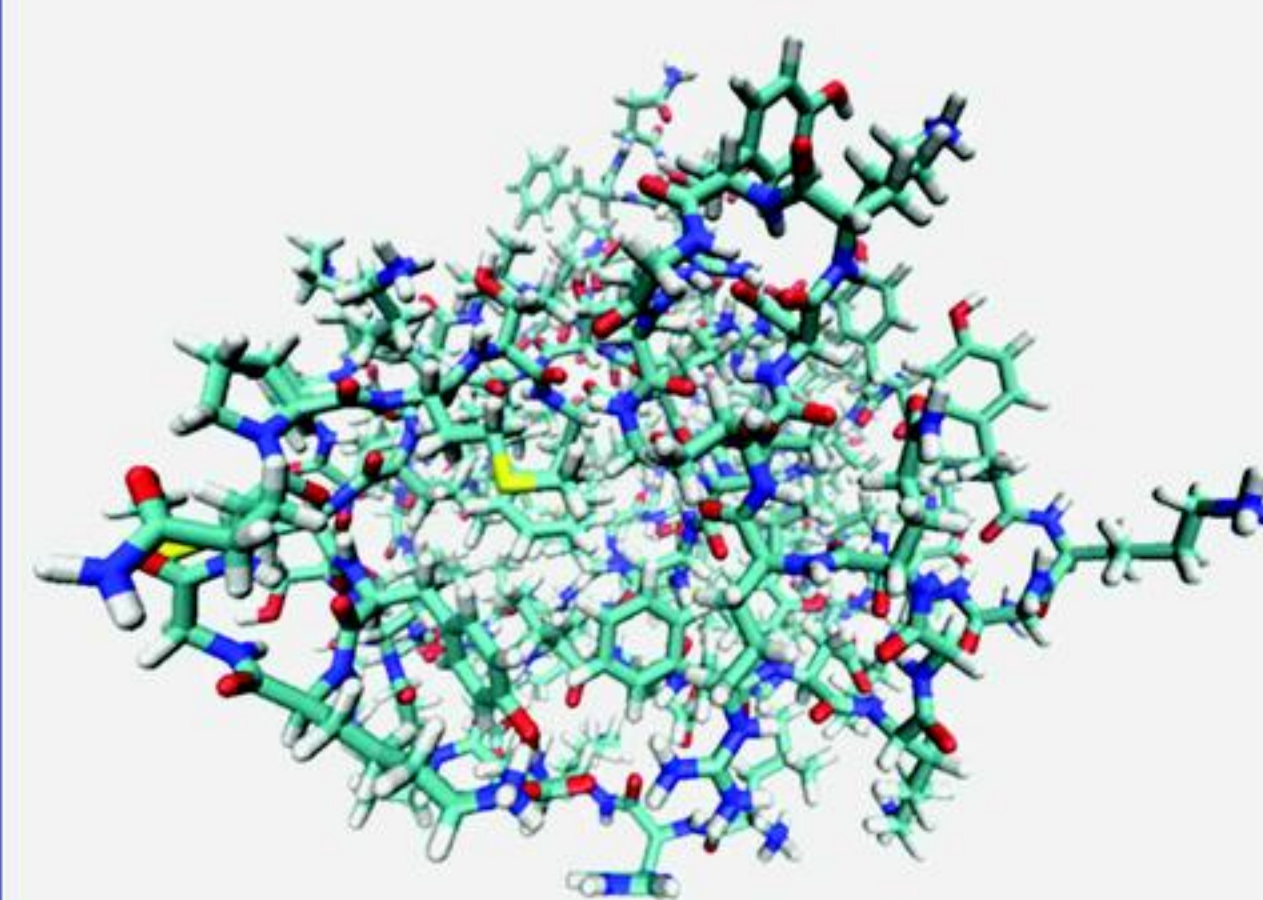
There are two benchmarks and four applications (one is a mystery) in the SC17 competition. We have spent great effort understanding and optimizing them. In particular, we have profiled the applications with Alinea, implemented optimizations in research papers, and studied checkpointing support to cope with power shutoff activities.

$$\begin{bmatrix} 1 & 1 & -4 \\ 0 & 1 & 7 \\ 0 & 0 & 79 \end{bmatrix}$$

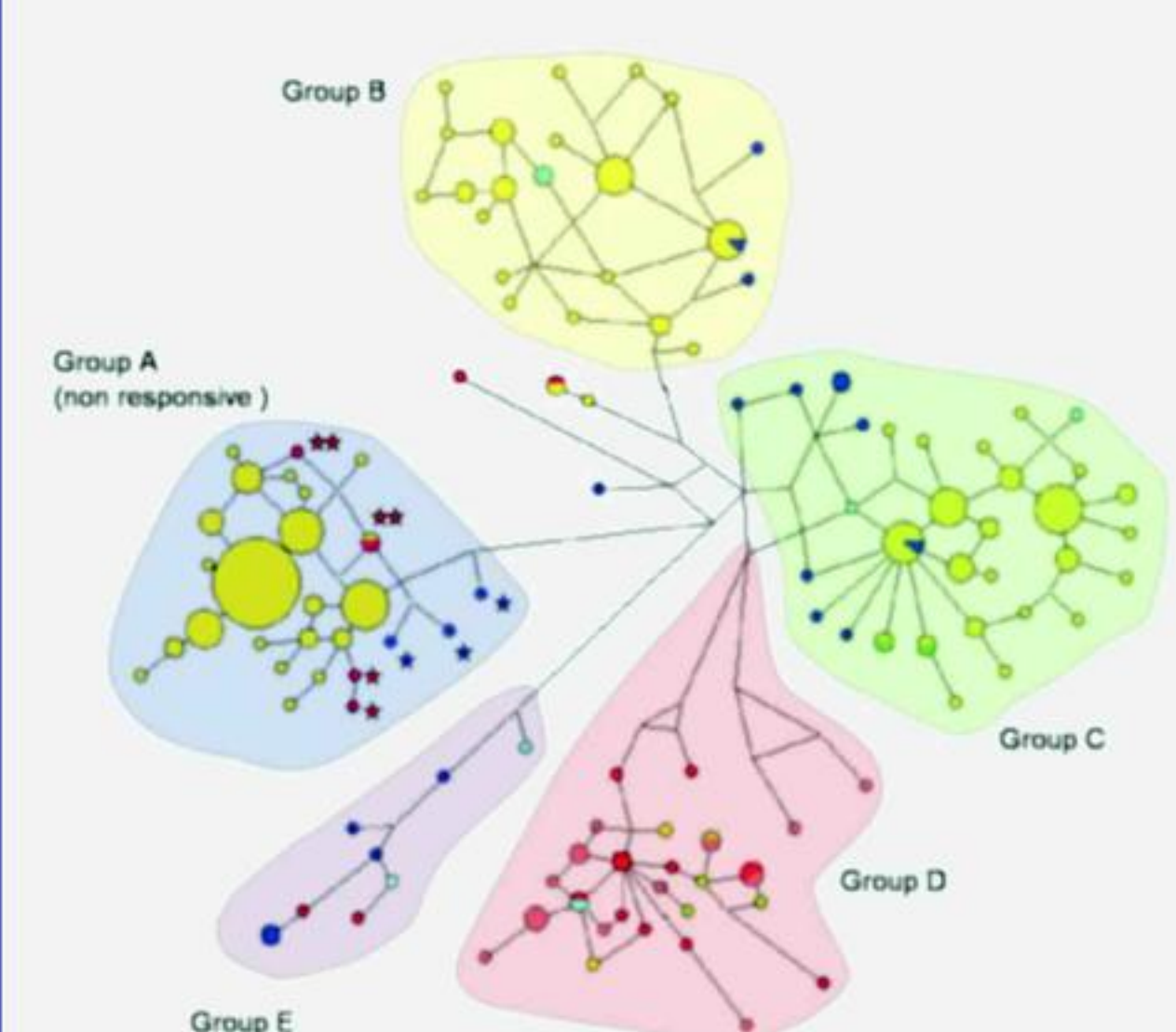


HPL & HPCG are two benchmarks used in the competition. HPL solves a dense system of linear equations with LU factorization while HPCG solves a sparse one with preconditioned conjugate gradient. We have carefully tuned both benchmarks on the Tesla V100s with power monitors and the latest binaries provided by NVIDIA.

Born is an application for seismic imaging using the reverse time migration technique. The core of the computation is a 3D finite difference kernel for modeling wave propagation. We found it is a memory-bound application through profiling. We have optimized the original CPU version with **loop tiling** [1], **data alignment**, **prefetching**, and **SIMD instructions**. We have also ported it to the GPU to utilize its high memory bandwidth and massive parallelism. After initial porting, we further optimized the GPU version with **data alignment** and **CUDA streams** [2].



LAMMPS is a parallel molecular dynamics simulation software supporting a variety of hardware architectures. In this competition, we are required to reproduce the performance of a portable Tersoff potential vectorization scheme which supports ARM, Intel CPU, Intel Xeon Phi, and Nvidia GPU. We have carefully studied the paper describing the algorithm and ran the experiments. We have also prepared scripts to extract and generate plots from the experiment data.



MrBayes is a parallel software that applies bayesian inference in phylogeny with a wide selection of evolution models. It uses the Metropolis Coupled Markov Chain Monte Carlo (MC3) algorithm. We have tuned a wide range of model settings to select the best one that fits our dataset. We found MrBayes to be compute-bound with profiling. Therefore, we optimized it with **FMA and AVX(-512) instructions** based on the original SSE version to speed up key compute-intensive functions.

Running Strategies

- Our goal is to maximize the utilization and efficiency of our cluster and the cloud component throughout the competition.
- Given the nature of application datasets, we plan to finish LAMMPS first and run Born, MrBayes, and the mystery application in the remaining time.
- As we have high performance GPUs on premise and AVX-512 supported CPUs in the cloud, we plan to run GPU-optimized applications (e.g. Born) on our cluster and CPU-optimized applications (e.g. MrBayes) on the cloud for **maximum performance-per-watt**.
- We will closely monitor the power consumption and adjust our workloads accordingly to prevent going over the power.

References

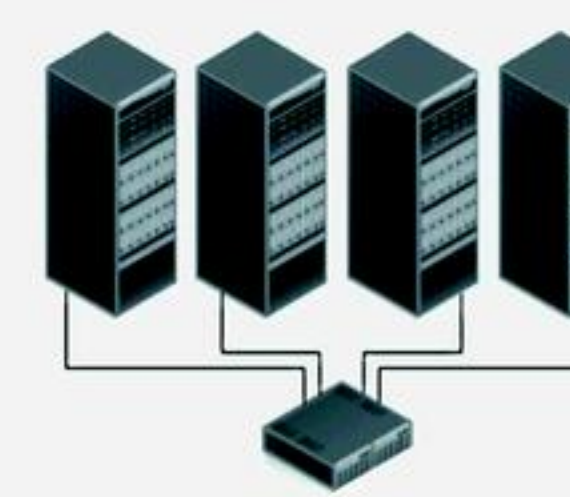
- [1] Rivera, G., & Tseng, C. W. (2000, November). Tiling optimizations for 3D scientific computations. In Supercomputing, ACM/IEEE 2000 Conference (pp. 32-32). IEEE.
[2] Micikevicius, P. (2009, March). 3D finite difference computation on GPUs using CUDA. In Proceedings of 2nd workshop on general purpose processing on graphics processing units (pp. 79-84). ACM.

Why We Will Win

Deep understanding of HPC knowledge



Powerful hardware



Put a lot of effort into optimizing the applications



Good team dynamics



Special Thanks

