

Constructing the Hyper-Kamiokande Computing Model in the Build Up to Data Taking

Sophie King^{1,*} on behalf of the Hyper-Kamiokande Collaboration

¹King's College London, UK

Abstract. Hyper-Kamiokande is a next-generation multi-purpose neutrino experiment with a primary focus on constraining CP-violation in the lepton sector. It features a diverse science programme that includes neutrino oscillation studies, astrophysics, neutrino cross-section measurements, and searches for physics beyond the standard model, such as proton decay. Building on its predecessor, Super-Kamiokande, the Hyper-Kamiokande far detector has a total volume approximately 5 times larger and is estimated to collect nearly 2 PB of data per year. The experiment will also include both on- and off-axis near detectors, including an Intermediate Water Cherenkov Detector. To manage the significant demands relating to the data from these detectors, and the associated Monte Carlo simulations for a range of physics studies, an efficient and scalable distributed computing model is essential. This model leverages Worldwide LHC Grid computing infrastructure and utilises the GridPP DIRAC instance for both workload management and for file cataloguing. In this report we forecast the computing requirements for the Hyper-K experiment, estimated to reach around 35 PB (per replica) and 8,700 CPU cores (~100,000 HS06) by 2036. We outline the resources, tools, and workflow in place to satisfy this demand.

1 Introduction

The Hyper-Kamiokande (Hyper-K) experiment [1, 2] is the successor to the highly successful and accomplished Super-Kamiokande (SK) [5, 6] and T2K (Tokai-to-Kamioka) [3, 4] experiments. Currently under construction, the Hyper-K far detector (HKFD) is a 258 kton underground water Cherenkov detector. The inner detector will house 20,000 inward-facing photomultiplier tubes (PMTs), 50 cm in diameter, along with a thousand multi-PMTs that each contain nineteen 8 cm PMTs. This region will be encompassed by the outer detector, which is designed to be instrumented with a few thousand 8 cm PMTs to veto incoming backgrounds. This amounts to tens of thousands of readout channels, and a post-trigger rate of around 5 TB/day of data to be transferred and stored off-site, setting the scale for Hyper-K storage requirements.

Hyper-K will search for CP-violation in the lepton sector through the study of neutrino oscillations in an accelerator-based long-baseline neutrino oscillation configuration. For this purpose, in addition to the far detector, Hyper-K features several near detectors, which will constrain systematic uncertainties relating to the flux and neutrino interaction models and measure the neutrino beam properties, as well as performing cross-section measurements

*e-mail: sophie.king@kcl.ac.uk

and searches for new physics. Hyper-K is also building the Intermediate Water Cherenkov Detector (IWCD), providing a near detector that utilises the same technology and target as the far detector, as well as the ability to profile the beam across a continuous range of off-axis angles.

The computing predictions that cover the raw and processed data needs of Hyper-K detectors, as well as the Monte Carlo (MC) production samples covering the signal, background and control samples needed for all physics analyses, are presented in Section 2. The infrastructure, tools and workflow that define the Hyper-K computing model, constructed to meet these needs in an organised and efficient manner, is defined in Section 3.

2 Hyper-K Computing Forecasts

The Hyper-K far detector is currently under construction, and the experiment will start taking data in 2027. The computing forecasts were split into two stages: the construction period, from 2023 up till the end of 2026¹; and the first ten years of operation, from the start of 2027 to the end of 2036². Internally within the Hyper-K Computing Group, a detailed year-by-year prediction has been forecast for each detector and for each sample type, with input from the Computing, Software, DAQ, Calibration and Physics working groups. However, due to the large uncertainties on these estimates - stemming from both the current phase of rapid software development within Hyper-K, which will stabilize over the next couple of years, and because aspects of the DAQ are still being finalised - we present only the 2023 and 2036 points, and use a straight line to convey the order of magnitude while avoiding the impression of year-to-year precision. The detailed breakdown is available upon request to the computing and funding agencies that support Hyper-K.

2.1 Aggregated CPU and Storage Projections

The Hyper-K computing requirements for raw data storage as well as both data and MC production, across all detectors, are totalled and presented as a function of time in Figure 1. The left axis indicates storage, which is projected to reach around 35 PB for a single replica of each file. The policy is to maintain three replicas of raw data, one in Japan and two in different countries, and two replicas of MC. The right axis represents the number of CPU cores and is predicted to reach nearly 8,700 cores ($\sim 100,000$ HS06).

Throughout the construction period, the computing estimates revolve solely around the MC simulations that are generated for physics sensitivity studies, detector design optimization, validation of trigger algorithms, and calibration studies. During this period, the IWCD MC production dominates the CPU demand and the HKFD MC dominates the storage.

During the operational phase, the storage and CPU associated with raw data storage, data (re)processing, calibration and event reduction/pre-selection contribute to the computing requirements. While the HKFD raw data quickly dominates the storage needs of the experiment, accounting for around 70% of the total, the IWCD MC production continues to dominate the CPU and is responsible for around 70% of the total.

¹It should be noted that construction began in 2021, but these estimates cover the period starting in 2023 as this is the point where our MC production campaigns are starting to ramp up to a more significant level.

²It is expected that Hyper-K will continue for an additional 10 years till the end of 2046. A *very* rough prediction for the end of this period can be considered by doubling the 2036 estimate.

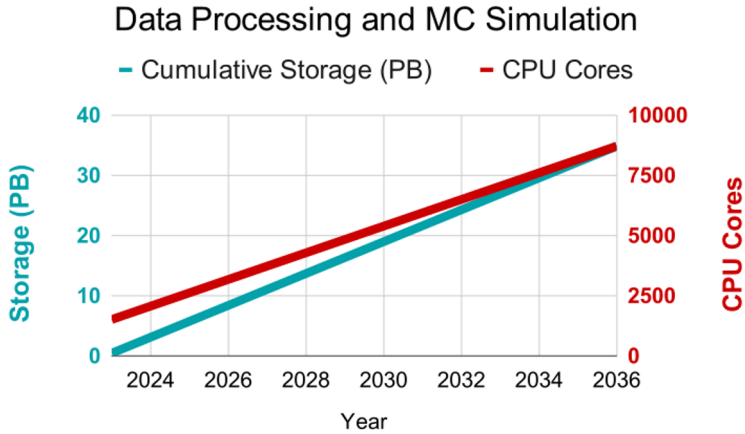


Figure 1: The total computing forecasts, covering all the Hyper-K detectors, considering both data and MC needs, and for all signal, background and control/calibration samples. Note that the storage estimate is per replica.

2.2 Resource-intensive tasks

In this section we comment on the tasks that drive the CPU and storage estimates during the operational phase of the experiment. This helps us to understand the best way to meet these demands and consider areas where we have potential to improve efficiency.

CPU usage for Water Cherenkov reconstruction: The CPU requirements of the IWCD MC production, which dominates both the construction and operational periods of the experiment, is due to the reconstruction stage. While the two Water Cherenkov detectors (HKFD, IWCD) use the same reconstruction algorithm, the close proximity of the IWCD to the neutrino beam results in a much larger number of beam events and hence requires MC samples with higher statistics. The current reconstruction algorithm, fitQun [7], was successfully introduced for SK to reduce the fiducial volume while also improving reconstruction capabilities, and is also used within both T2K and Hyper-K. It uses a maximum-likelihood based algorithm that computes the likelihood for different particle topologies along with their reconstructed kinematics. Efforts are being made to improve reconstruction processing times, both by improving the efficiency of the existing method, and by investigating machine learning based techniques, which seek to improve not only speed but also reconstruction capabilities [8, 9].

Raw data storage: The vast size of the far detector, coupled with tens of thousands of readout channels, yields a predicted (post-trigger) raw data rate of nearly 2 PB/year to be transferred and stored off-site. While this estimate assumes a basic level of data rejection due to triggering, it is a slightly conservative number and has the potential to decrease once the triggering is finalised and the collaboration reaches agreement about what data to store long-term. The majority of this data will be archived to tape and brought onto disk upon request; regularly accessed data samples, based on triggers and reduction/pre-selection cuts, will be replicated to disk storage for continued ease of access.

2.3 Future Projections

As previously mentioned, machine learning based techniques are under investigation to improve both processing speeds and the performance of kinematic reconstruction and particle identification. Other stages of production for which machine learning methods are being explored include simulation and calibration, albeit with lower priority. These developments could significantly alter the Hyper-K computing needs, introducing a need for GPU-based jobs within the computing workflow and resource allocations. These innovations could drastically reduce processing times, though it is expected that some level of likelihood-based processing would remain in conjunction with the new methods, at least for validation purposes. While these considerations are under internal discussion, they are not included in the current computing projections. As the timeline and performance metrics become more refined, we plan to update the computing forecasts accordingly.

The current computing estimates focus solely on the needs concerning the storage and production of real data and MC simulations. Internal discussions between the computing and analysis groups are underway to understand the computing requirements of the different high-level analysis frameworks. As neutrino physics enters a systematically dominated era, statistical analysis in increasingly high-dimensional space can no longer be performed on small institute clusters. Physics groups are instead turning to HPC clusters, with many of these analysis frameworks utilising GPU acceleration. As the data grows in size and the complexity of these analyses increases, resource limitations can lead to the need to introduce approximations that improve speed but decrease accuracy, or to compromise on which studies can be performed. Preparing results for significant conferences can also impose strict deadlines, resulting in large spikes in demand. Integrating analysis level tasks into the computing framework can help to prioritise at the expense of less urgent tasks, and including the computing requirements in negotiations with computing and funding bodies can help to better plan for and guarantee these needs. In future iterations, the computing projections and workflows for computationally demanding analysis will be presented alongside the needs of real data and MC simulation.

3 The Hyper-K Computing model

Hyper-K adopts a tiered system for computing, similar to that of the Worldwide LHC Computing Grid (WLCG), and benefits from much of its infrastructure and tools. To meet specific requirements of Hyper-K jobs and data management, we utilize community-based tools which we integrate with custom Hyper-K software. This is discussed in Section 3.1, while the tier definitions are outlined Section 3.2.

3.1 Computing Tools, Production Workflow and Data Management

3.1.1 Distributed Computing with GridPP DIRAC

The DIRAC (Distributed Infrastructure with Remote Agent Control) project [10] is a software framework that interfaces users with computing resources. It offers a pilot-based Workload Management System (WMS) for job submission, which can connect to grid, cloud and local batch systems. DIRAC also provides a Data Management System (DMS) for cataloging file replication and metadata, in the Dirac File Catalogue (DFC), along with the associated end-user tools to manage this. Hyper-K uses both of these services, provided by the

GridPP instance of DIRAC hosted at Imperial College London [12]. This is a multi-virtual organisation (VO) service³, with Hyper-K being connected through the hyperk.org VO.

All T1 and T2 Hyper-K storage is managed through the DIRAC DMS and DFC. Efficient data transfer between sites is possible through third party services such as FTS3 [13]. For bulk data operations Hyper-K uses FTS3 with the Dirac Request Management System (RMS), which provides a scheduling and monitoring service.

At the time of writing, Hyper-K uses DIRAC for job submission to access both allocated and opportunistic resources at UK (RAL, ICL and other T2 sites), Italian (INFN) and French (IN2P3) grid sites. Through the JENNIFER2 project (funded by the Horizon 2020 programme⁴) and collaboration between the T2K, Hyper-K and Belle II [14] experiments, Hyper-K connected cloud resources hosted by INFN and IN2P3 to the GridPP DIRAC instance using the virtual machine manager VCYCLE [15] to generate DIRAC pilot based VMs. Following from this work, a cloud demonstrator was deployed that integrates token based authentication in the Openstack module of VCYCLE, allowing European Grid Infrastructure (EGI) cloud resources [16] to be accessed by DIRAC [17]. This demonstrates versatile nature of DIRAC, facilitating the integration of cloud resources into the central pool for Hyper-K resources.

3.1.2 Software Distribution

All Hyper-K production is done with version controlled containers, which ensures consistent and reproducible results. Hyper-K uses the CERN Virtual Machine File System (CVMFS) [18] configured for EGI and hosted by RAL (UK), where these containers are distributed as singularity sandboxes.

3.1.3 Hyper-K Computing tools

The Hyper-K computing tools are written in python and designed to be platform-independent, such that a coherent set of tools and job definitions can be used across multiple resources, with a different wrapper written for each type of backend. For submission to DIRAC-linked sites, the tools utilize the python-based DIRAC API. Similarly, for all DFC data management, the Hyper-K tools integrate the DIRAC API to make use of the DMS and RMS tools for both single and bulk file operations and metadata management. The DIRAC software may be installed locally, accessed on CVMFS or containerised with the Hyper-K tools, such that the minimum requirement from a site that wishes to use the full extent of the Hyper-K computing package is that it has open the necessary ports to access the GridPP DIRAC server⁵.

3.2 Computing infrastructure and Tiers

3.2.1 Grid/DIRAC Tiers

This section outlines the Hyper-K tiers that are connected to the hyperk.org VO through the GridPP instance of DIRAC. While this is primarily grid resources, it may also include cloud and batch systems, as mentioned in Section 3.1.1.

³In grid computing, a virtual organization allows a set of individuals or institutions to securely share computing resources

⁴The JENNIFER2 project is funded under the Horizon2020 program of the European Union as a Marie Skłodowska Curie Action of the RISE program: MSCA-RISE-2018 call, project JENNIFER2, GA 822070

⁵Though note that access to the DIRAC server is not necessary to use all aspects of the Hyper-K computing tools.

Tier-0 (T0) sites are where raw data from the detectors is initially archived onto tape. This is KEK Central Computer System (KEKCC) for the near detectors, and a dedicated computing system at the Kamioka Observatory for the far detector; both are based in Japan, close to the associated detector sites. Raw data processing of each detector should be performed at the corresponding T0 site, although DIRAC-linked resources may be used as an over-spill if required, e.g. if a large amount of data needs to be reprocessed at short notice. It is planned that 3,000 cores will be dedicated to this purpose at the Kamioka Observatory. This is around double what is estimated for real-time processing of the HKFD, to allow for any reprocessing needs. For the near detectors, the CPU required for data processing is a small fraction of their MC production CPU needs. However, to ensure data processing and reprocessing can be performed in a timely manner and at short notice, the number of cores allocated is based on being able to reprocess a years worth of beam data in around 2 months. This will initially be set at around 1,000 cores, with the requirement increasing as more data is collected. When these machines are not in use for data processing, these resources can be allocated for other purposes such as calibration or MC production.

Tier-1 (T1) sites hold copies of the raw and processed data, as well as MC production batches, on a mix of disk and tape. They will also generate and process most of the MC simulations. The countries that are Tier-1 during the construction phase, and intend to maintain this status throughout operation, are France, Italy and the UK.

Tier-2 (T2) sites provide disk storage that is mostly utilised in a temporary manner based on demand. T2 sites can support the MC production campaigns with CPU, or may be used for specific tasks, especially in instances where the disk is utilized such that jobs can access the files locally. Some or all of these countries that are T1 will also provide resources at T2 sites. Other countries that are looking to secure grid resources for Hyper-K are Canada, Japan, Poland, Sweden and Switzerland.

3.2.2 *Standalone Tiers*

The Hyper-K computing model is not confined to DIRAC-linked resources; some countries will contribute ‘standalone’ computing resources such as local HTC and HPC clusters, cloud resources, and client-server storage solutions. It is worth nothing that local clusters have the potential to be connected to DIRAC via SSH, but most will likely be treated as standalone. Cloud resources may also be added to DIRAC; whether Hyper-K connects them to DIRAC or treats them as a standalone will depend on both the purpose of these resources and any technical considerations for a given site.

Standalone Compute-1 (SC1) sites provides CPU/GPU at similar level of service as as T1; an example is the Kamioka Observatory computing system (which will provide both SC1-CPU and T0-Storage services), or the resources allocated to Hyper-K through The Digital Research Alliance of Canada.

Standalone Compute-2 (SC2) resources are defined as a site that commits to providing the resources and person power to manage a specific computing service task; these are typically institute clusters, and the Hyper-K computing group does not interact directly with these resources, only with the collaborator assigned to the task.

Standalone Storage-3 (SS3) storage is any non-DIRAC/non-grid storage that is accessible to all Hyper-K collaborators. It may be used for analysis files or general-purpose file sharing. The KCL (UK) hosted Nextcloud service is an example of this.

4 Summary

The Hyper-Kamiokande experiment has a rich and diverse physics programme, enabled by its suite of specialised, yet multi-purpose, detectors. This gives rise to substantial computational and storage demands, met by significant infrastructure and pledges contributed by France, Italy, Japan, and the UK. To coherently manage these contributions, and to incorporate a heterogeneous pool of resources from additional sites and platforms going forward, a centralised, efficient, and well-coordinated computing model is essential. This is supported and realised by a combination of community software and services, notably DIRAC and the services provided by GridPP, and dedicated Hyper-K computing tools. As a result, the Hyper-K computing group can enable measurements based on the latest data to be published in an efficient and timely manner, while ensuring the long-term preservation of data and metadata.

5 Acknowledgements

The author would like to thank the computing services and support provided by the following. The Digital Research Alliance of Canada, GridPP, INFN-CNAF, the Centre de Calcul de l'IN2P3, KEKCC, and Kamioka Observatory (ICRR, The University of Tokyo) for computing resources and services. The European Union's Horizon 2020 Research and Innovation Programme for funding the JENNIFER2 project that supported part of this work. The Belle II collaboration, in particular S. Pardi, for collaboration on the cloud demonstrator. The GridPP Collaboration and Imperial College London for the GridPP DIRAC services, in particular the support from D. Bauer and S. Fayer. The CVMFS taskforce at UKRI RAL, funded by EGI.

References

- [1] Hyper-Kamiokande Collaboration, arXiv:1805.04163 [physics.ins-det] (2018)
- [2] Hyper-Kamiokande Collaboration, arXiv:2009.00794 [physics.ins-det] (2021)
- [3] T2K Collaboration, Nucl.Instrum.Meth.A **659** 106-135 (2011)
- [4] T2K Collaboration, Nature **580**, 339 (2020)
- [5] Super-Kamiokande Collaboration, Nucl. Instrum. Meth. A **501** 418-462 (2003)
- [6] Super-Kamiokande Collaboration, Phys. Rev. Lett. **81** 1562-1567 (1998)
- [7] Super-Kamiokande Collaboration, Prog. Theor. Exp. Phys. 2019, **053** F01 (2019)
- [8] B. Jamieson et al., Front.Big Data **5** 978857 (2022)
- [9] N. Prouse, PoS ICHEP2020 919 (2021)
- [10] A. Tsaregorodtsev et al., J. Phys. Conf. Ser. **119** 062048 (2008).
- [11] The GridPP Collaboration, J. Phys. G **32** N1-N20 (2006)
- [12] Bauer D., Fayer S., J. Phys.: Conf. Ser. **898** 052003 (2017)
- [13] A. A. Ayllon et al, J. Phys.: Conf. Ser. **513** 032081 (2014)
- [14] The Belle II Collaboration, arXiv:1011.0352 [physics.ins-det] (2010)
- [15] McNab A., Love P., MacMahon, E., J. Phys.: Conf. Ser. **664** 022031 (2015)
- [16] The European Grid Infrastructure. Retrieved from <http://www.egi.eu>
- [17] S. Pardi, S. King, M. Guigue, et al., Int. J. Appl. Phys. Math. **13** 1 (2023)
- [18] P. Buncic et al., J. Phys.: Conf. Ser. **219** 042003 (2010)