

Kubernetes for AI-Enabled Scientific Research Computing, and Education

PEARC 26 Tutorial, July 27, 9AM-5PM



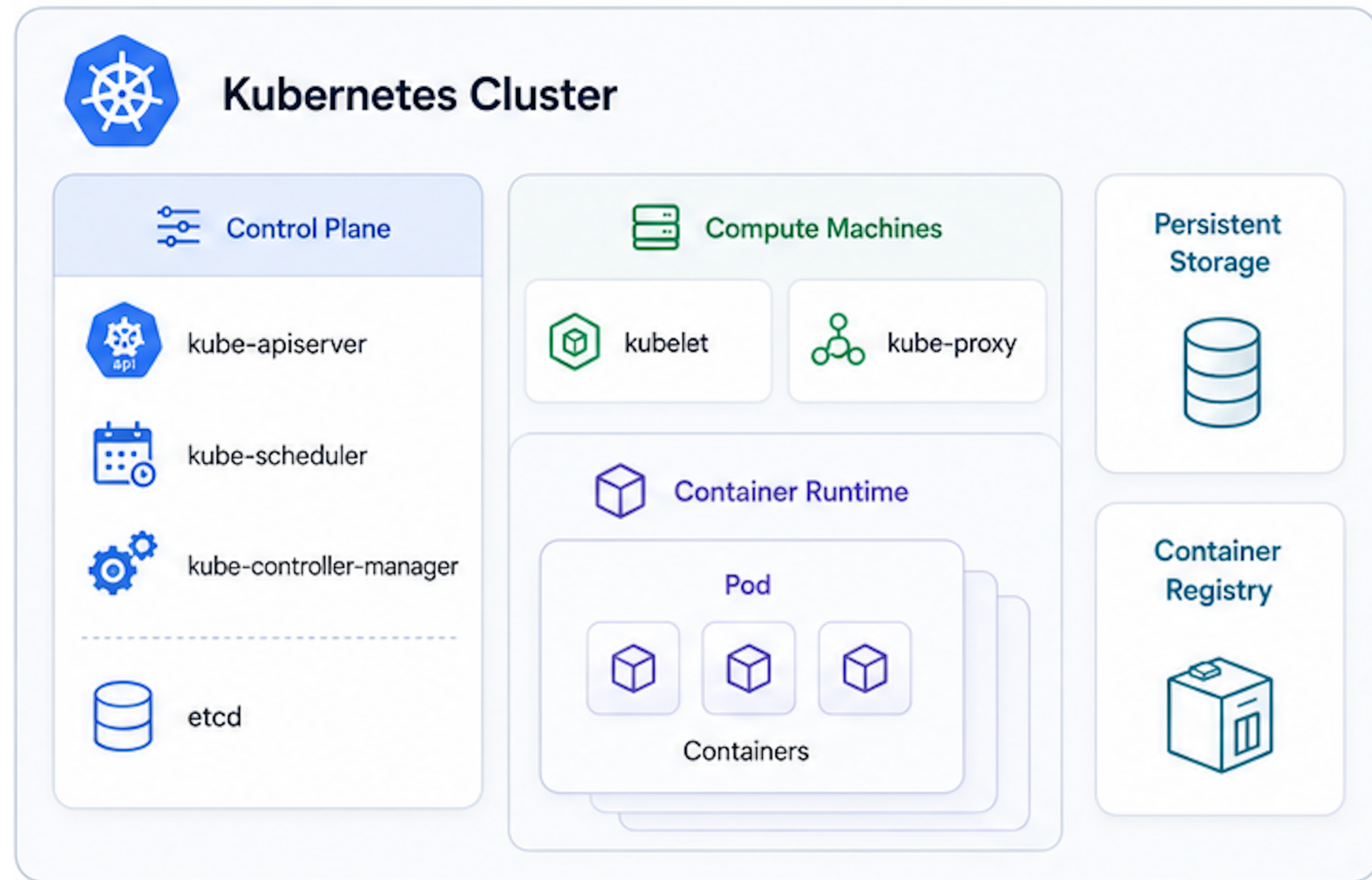
SAN DIEGO
SUPERCOMPUTER CENTER

Agenda

- Welcome, Kubernetes & NRP Architecture
- Basic Docker & Kubernetes Hands-On
- Coffee Break (10:30-11:00AM)
- AI & Computational Science Applications
- Persistent Storage & I/O for AI/Scientific Workloads
- Lunch Break (12:30-1:30PM)
- JupyterHub on NRP
- Advanced — Deploy a Custom JupyterHub & Build Images in NRP GitLab
- Coffee Break (3:00-3:30PM)
- Advanced — Deploy a Custom JupyterHub & Build Images in NRP GitLab (contd)
- Get your own NRP access + wrap-up / Q&A

What is Kubernetes (K8s)?

- Kubernetes is a platform for running **containerized applications** across many machines.
- The **control plane** manages the cluster and decides where workloads run.
- **Worker nodes** run the actual applications inside pods.
- Kubernetes automatically handles scheduling, restarts, scaling, networking, and storage.
- Users request resources such as CPU, memory, GPUs, and storage instead of managing individual machines.
- **Nautilus** uses Kubernetes to support large-scale research computing, services, and specialized hardware.



K8s = K u b e r n e t e s
1 2 3 4 5 6 7 8

Docker/Containers



- **Docker** is a platform for building, packaging, and running applications in **containers**.
 - **Containers** are lightweight, portable, and consistent environments.
 - They include everything your app needs: code, libraries, and dependencies.
 - Docker makes it easy to run the same app on your laptop, a server, or in the cloud.
- Why Docker matters for Kubernetes:
 - Kubernetes doesn't build containers, it runs them.
 - Docker (or other container runtimes) provides the container images that Kubernetes manages.
- NRP hosts its own **container registry** with gitlab

Kubernetes Resources

- **Pod** - smallest deployable unit in K8s. Consists of one or more containers. Ephemeral
- **Job** - Process that runs until a desired condition is met. Can launch one or multiple pods
- **Deployment** - Persistent process; complex orchestrations e.g., JupyterHub
- Persistent Volume Claim (PVC) - Represents a request for storage. PVC is a namespace scoped object, storage is cluster scoped.
 - Keep in mind that pods are ephemeral workspaces. PVCs allow you to save your work

PEOPLE & ACCESS



Cluster Admins

Manage the entire cluster and all namespaces



Namespace Admins

Manage resources within assigned namespaces



Users / Developers

Create and manage resources in their namespace



View Only

Can view resources based on permissions



Kubernetes Cluster



Cluster-scoped resources shared across all namespaces



Nodes



StorageClasses



PersistentVolumes



ClusterRoles



CustomResourceDefinitions



Namespace: research



Namespace Admin

NAMESPACE-SCOPED RESOURCES



Deployments



Services



ConfigMaps



Secrets



Roles & RoleBindings



Pods



Pods



Namespace: data



Namespace Admin

NAMESPACE-SCOPED RESOURCES



Deployments



Services



ConfigMaps



Secrets



Roles & RoleBindings



Pods



Pods



Namespace: ml



Namespace Admin

NAMESPACE-SCOPED RESOURCES



Deployments



Services



ConfigMaps



Secrets



Roles & RoleBindings



Pods



Pods



RESOURCE SCOPE



Namespace-scoped

Exist within a specific namespace.
Isolated between teams.



Deployments



Services



ConfigMaps



Secrets



Roles & RoleBindings



Pods



Cluster-scoped

Exist across the entire cluster.
Shared globally.



Nodes



PersistentVolumes



StorageClasses



ClusterRoles

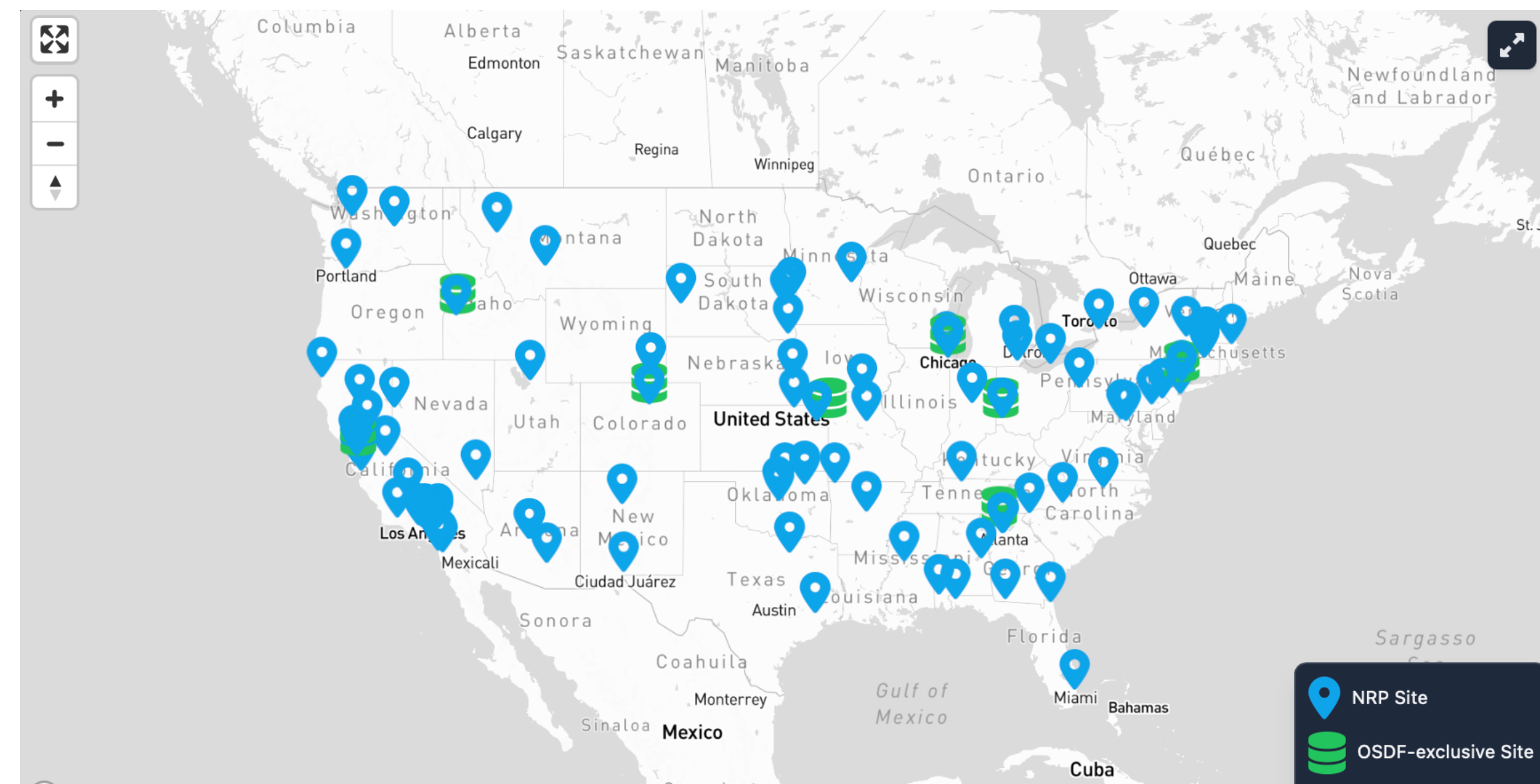


CustomResourceDefinitions

What is NRP?

- Internationally distributed
 - Most sites are in the US
- Scale of NRP
 - ~500 compute nodes
 - ~32K CPU cores
 - ~1400 GPUs

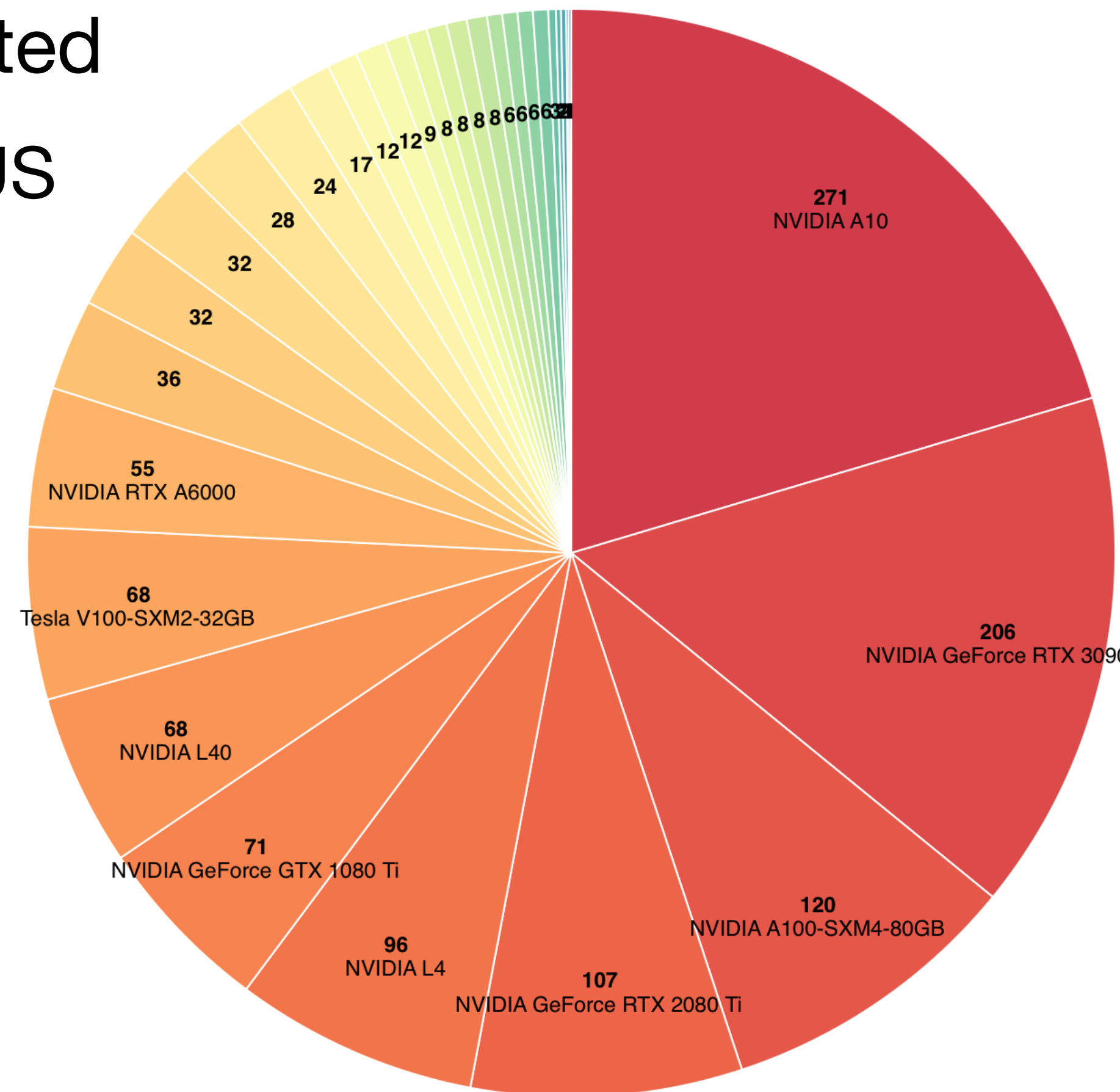
<https://dash.nrp-nautilus.io/>



NRP is a shared national cyberinfrastructure built on the Nautilus Kubernetes cluster. It provides hundreds of GPU nodes, shared storage, and managed services such as JupyterHub, S3, a vector database (Milvus), and an OpenAI-compatible managed LLM endpoint — all free for U.S. academic research, teaching, and outreach.

What is NRP?

- Internationally distributed
 - Most sites are in the US
 - Scale of NRP
 - ~500 compute nodes
 - ~32K CPU cores
 - ~1400 GPUs
- 

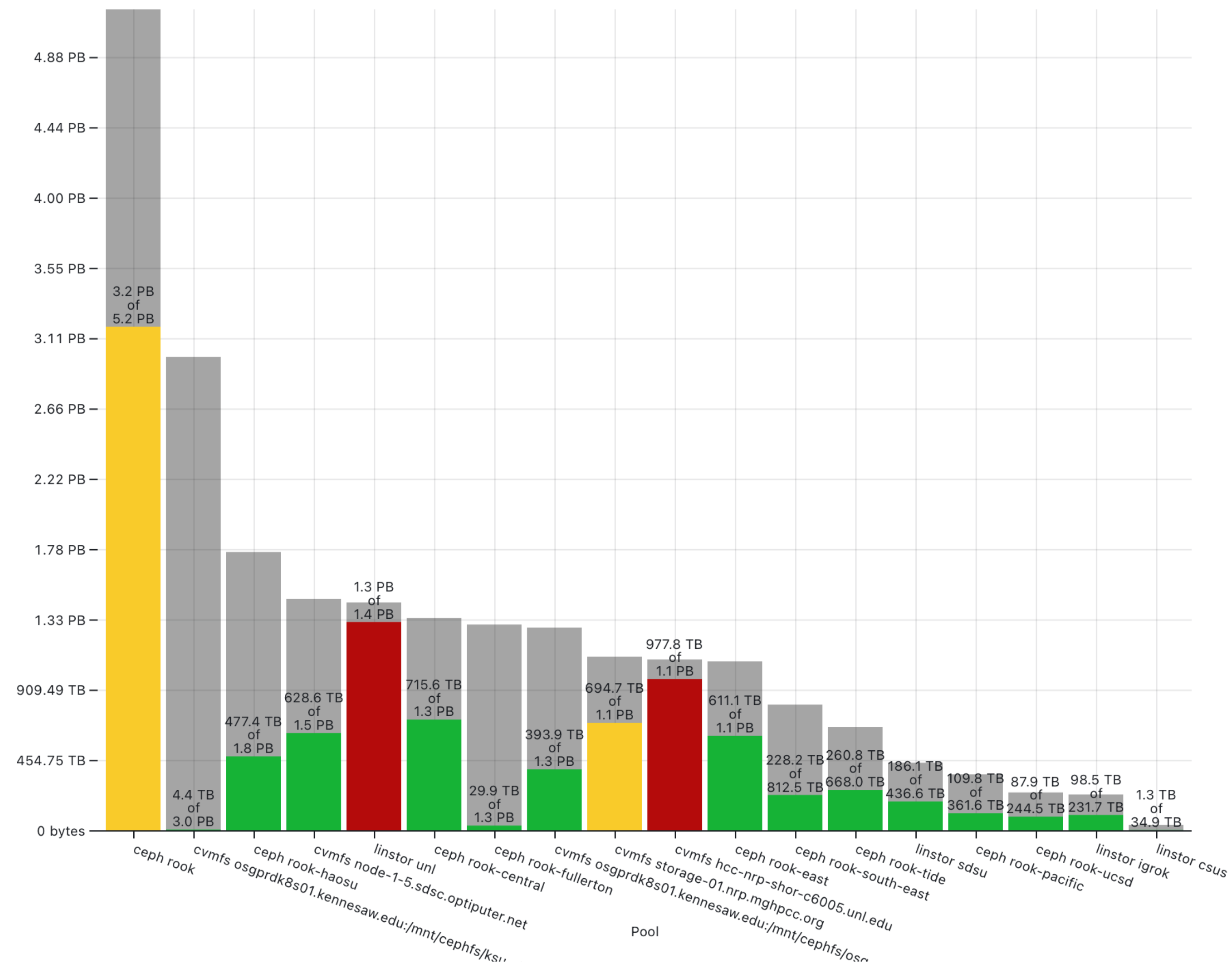


GPU Table

GPU Model	Number
NVIDIA A10	271
NVIDIA GeForce RTX 3090	206
NVIDIA A100-SXM4-80GB	120
NVIDIA GeForce RTX 2080 Ti	107
NVIDIA L4	96
NVIDIA GeForce GTX 1080 Ti	71
NVIDIA L40	68
Tesla V100-SXM2-32GB	68
NVIDIA RTX A6000	55
NVIDIA A100 80GB PCIe	36
Tesla V100-SXM2-16GB	32
NVIDIA RTX A4000	32
NVIDIA A100 80GB PCIe	28
NVIDIA GeForce RTX 4090	24
NVIDIA TITAN Xp	17
NVIDIA L40S	12
NVIDIA A40	12
NVIDIA A100-PCIE-40GB	9
NVIDIA GeForce GTX 1080	8
NVIDIA RTX 5000 Ada Generation	8
NVIDIA RTX PRO 6000 Blackwell Max-Q Workstation Edition	8
NVIDIA H100 80GB HBM3	8
NVIDIA H200 NVL	6
Tesla V100-PCIE-16GB	6
NVIDIA TITAN RTX	6
NVIDIA RTX A5000	6
Tesla T4	3
NVIDIA RTX 4000 Ada Generation	2
Quadro RTX 6000	2
Quadro RTX 8000	1
NVIDIA TITAN X (Pascal)	1

Storage Options

- Ceph File System → **Ceph FS** → RWX
- Ceph Rados Block Device → **Ceph RBD** → RWO
- **Linstor** Block → RWO
- **S3** Compatible storage → Bucket Endpoints
- Open Science Data Federation → **OSDF**
- **No Archival Storage**
 - Stale data will be deleted after 6 months if not used.



Access Methods



Kubernetes



JupyterHub



Coder

- Kubernetes: Write yaml files and use `kubectl` command line tool.
- Jupyterhub*: <https://jupyterhub-west.nrp-nautilus.io>
- Coder*: <https://coder.nrp-nautilus.io>

***Note:** These are the NRP provided JupyterHub and coder, there are others.
Share the url you use to access these services when asking for help

JupyterHub on Nautilus

- <https://jupyterhub-west.nrp-nautilus.io/>
- <https://jupyterhub-east.nrp-nautilus.io/>
- We host two JupyterHub instances
- Login via CILogon with your university credentials
- You can choose resources like CPUs, GPUs, GPU type, FPGAs, RAM, etc
- NRP provides several pre-built images
- Details documentation here:
<https://nrp.ai/documentation/userdocs/jupyter/jupyterhub-service/>

Server Options

/home/jovyan is persistent volume, 5GB by default. Make sure you don't fill it up - jupyter won't start next time. You can ask admins to increase the size.

The storage is created in West ceph pool by default. You can ask admins to move it to a different region.

[Available resources page](#)

[GPU types guide](#)

Contact admins in [Matrix](#).

Region

Any

GPUs

0

Cores

1

RAM, GB

8

FPGAs

0

GPU type

Any general



JupyterHubs on Nautilus Cluster

- PNRP Staff maintained and supported two JupyterHub instances
 - Institutional credentials to login using CILogon
 - Provide several container images
 - A total of 4942 users of Jupyter across 161 JupyterHubs in the past 12 months. This was 84% of NRP users.
- Institutions/Research Groups have cloned and customized our JupyterHub instances
 - Total of 161 JupyterHubs
 - Admins share information and recipes via gitlab, Slack, and Matrix channels

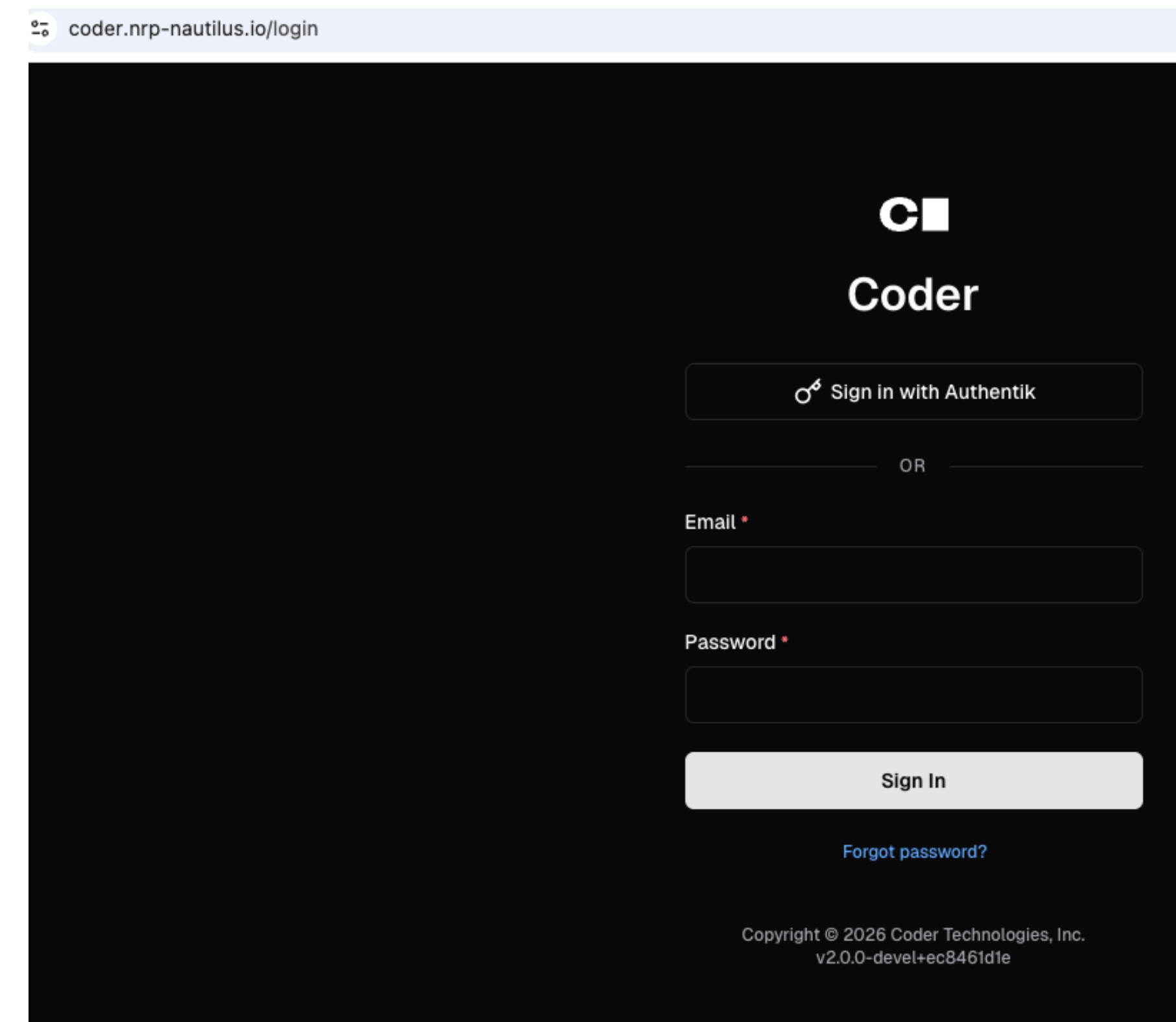
JupyterHub Images: NRP's Images

<https://nrp.ai/documentation/userdocs/running/sci-img/>

- NRP Deep Learning & Data Science Full, PyTorch
- NRP Deep Learning & Data Science Full, TensorFlow
- Selkies Desktop (Experimental) (ubuntu:mypasswd)
- NRP R Studio Server
- NRP Matlab
- NRP v1.3.1 (deprecated)
- NRP Desktop GUI
- NRP Desktop GUI + Relion
- NRP Desktop GUI + PRISM
- NRP OSGEO
- NRP FPGA
- NRP test

Coder on Nautilus

- <https://coder.nrp-nautilus.io/>
- We host a Coder instances
- Integrated with Authentik and allows login via CILogon with your university credentials
- You can choose resources like CPUs, GPUs, GPU type, FPGAs, RAM, etc
- NRP provides multiple templates
- More documentation here:
<https://nrp.ai/documentation/userdocs/coder/coder/>



Coder on Nautilus

- **Coder** provides a quick-start development environment similar to JupyterHub, running directly on our cluster. No Kubernetes knowledge required!
- **FPGA Requests:** Easily request and use Xilinx Alveo U55C FPGAs for specialized workloads
- **Preconfigured with Xilinx Vivado/Vitis License Server**, including **P4 license support**, so you can compile and deploy a variety of FPGA workloads right out of the box
- **Git Integration:** All workspaces come with SSH keys for seamless GitLab connections
- **User guide** for deploying custom Coder instance on Nautilus
<https://nrp.ai/documentation/userdocs/coder/deploy/>

LLM Support on NRP

- NRP provides several hosted open-weights LLM for either API access or use with our hosted chat interfaces.
- Chat interfaces – NRP Open WebUI (all NRP-hosted models), Cherry Studio (desktop application), and Chatbox (desktop application)
- LLM API access using Envoy AI Gateway – users can create tokens once they are part of LLM enabled namespace.
- Several models supported on NRP: qwen3, qwen3-small, gpt-oss, gemma, and qwen3-embedding. Models continuously evaluated and added to generally supported list if appropriate
- Models being evaluated: gemma-small, kimi, glm-5, and minimax-m2
- Documentation:

<https://nrp.ai/documentation/userdocs/ai/llm-managed/models/>

Move to hands on!

- The following link was emailed to attendees who registered for tutorial

<https://training.nrp-nautilus.io/pearc26/index.html>