

GPU PaaS

Reference Architecture for Cloud Providers and Enterprise Private Clouds

TABLE OF CONTENTS

Introduction	3
Delivering GPU PaaS capabilities with the Rafay Platform	3
Operating Model	8
Rafay Platform Deployment	13
Prerequisites for SaaS Consumption	13
Network Requirements	13
Security Settings	13
Prerequisites for Self-Hosted Controller Deployment	14
Operating System	14
Hardware Requirements	14
Network Requirements	14
Security Settings	14
Other Considerations	14
Additional Prerequisites	14
System Support	15
Deployment & Management	16
NVIDIA-Certified System Hardware and Network Topology	16
Rafay Platform Configuration	18
Cloud Provider SKU Management	18
Kubernetes Cluster Reference SKU Template	19
Virtual Cluster Reference SKU	20
NVIDIA Services Deployment Configuration	21
NVIDIA NIM Deployment Reference Template	22
NVCF Deployment Reference Template	23
NVIDIA Run:ai Deployment Reference Template	24
AI Applications Developed by Cloud Providers	25
Troubleshooting and Support	25
Troubleshooting	25
RAFAY Support	25
Example Use Cases	26
Supporting Documentation	27

Introduction

This document furnishes the reference architecture for cloud providers leveraging NVIDIA's accelerating compute infrastructure and the Rafay Platform to deliver a Platform-as-a-Service (PaaS) experience that enables developers and data scientists to consume GPUs on demand. The Rafay Platform also empowers cloud providers to include AI applications developed by NVIDIA and other partners as part of the GPU PaaS experience.

The Rafay Platform delivers capabilities through multiple programmatic interfaces, so cloud providers can operate an NVIDIA GPU PaaS with ease. Cloud providers can consume the Rafay Platform as SaaS or deploy the Rafay Platform in their self-hosted environments. In both cases, cloud providers can either white label the Rafay Platform, or can leverage the Platform's APIs to embed Platform capabilities into their pre-existing product portal(s).

The reference architecture furnished in this document can also be leveraged by cloud providers that meet the [NVIDIA Cloud Partner](#) (NCP) designation requirements as set forth by NVIDIA, as well as by enterprises building out NVIDIA GPU Clouds for internal consumption. The guide does assume NVIDIA GPU, Networking, and NVIDIA AI Enterprise software.

Delivering GPU PaaS capabilities with the Rafay Platform

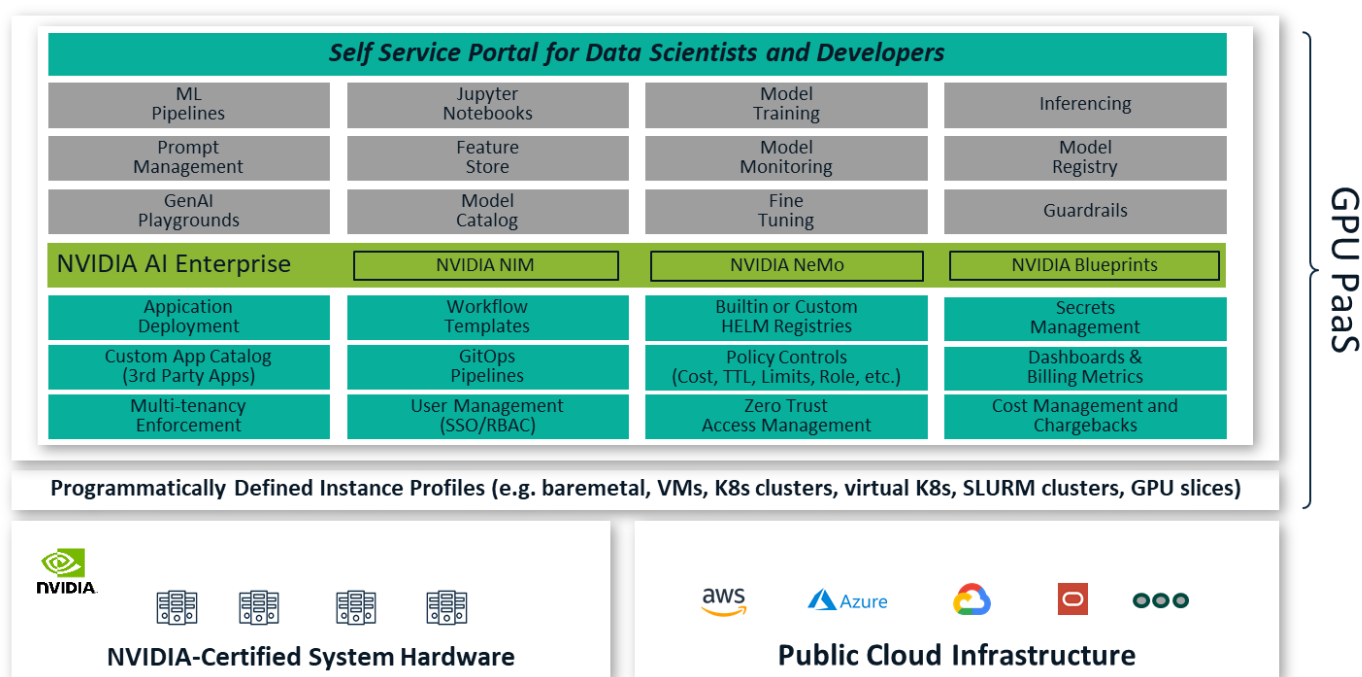
The Rafay Platform enables cloud providers to deliver GPU and CPU resources, AI & Generative AI services, and associated platform tools to their customers on top of NVIDIA's accelerated computing infrastructure. Through the Rafay Platform, cloud providers can deliver a curated GPU PaaS experience to their customer base such that the cloud provider is able to offer different experiences to different customers depending on need.

- The Rafay Platform allows cloud providers to allocate the GPU infrastructure to their customers in a self-service manner with multiple customers through several capabilities.
- The Rafay Platform enables cloud providers to programmatically offer various flavors and sizes of GPU and compute resources such as bare metal nodes, Kubernetes clusters, virtual Kubernetes clusters with multiple or fractional GPUs, SLURM clusters, etc.

- The Rafay Platform supports out-of-the-box integrations with NVIDIA services such as NVIDIA NIM™, or services such as Run:ai, etc., empowering cloud providers to offer NVIDIA's software solutions to their customers in a turnkey fashion.

Apart from the capabilities above, capabilities such as enterprise single sign-on (SSO), customer chargeback, cost management, policy enforcement, application deployment pipelines, visibility and persona-specific dashboard, etc., are also provided by the Rafay Platform. For the cloud provider's platform engineering teams, capabilities such as tenant administration and ease of troubleshooting for support teams, which are critical for deploying and operating a GPU PaaS, are also provided.

Rafay Platform provides cloud providers a turnkey GPU PaaS solution that comes packaged with critical capabilities highlighted in the below diagram to accelerate GPU Cloud deployments.



PaaS Reference Architecture for GPU Clouds

As shown above, NCPs can leverage the Rafay Platform in their data centers alongside their NVIDIA accelerated computing deployments, or they can leverage the Rafay Platform in public clouds to augment GPU capacity. The Rafay Platform will ensure that the PaaS experience is consistent across all environments, delivering a homogenous, single-pane-of-glass experience to developers and data scientists.

With the Rafay Platform, NCPs can deliver various compute packages, such as Kubernetes clusters, virtual Kubernetes clusters, SLURM clusters, fractional GPUs, etc., to their customers. The Rafay Platform provides foundational capabilities such as SSO/RBAC, zero trust access, GitOps pipelines, policy enforcement, persona-specific dashboards, and several other critical capabilities that enterprises will expect from NCPs to adopt the NCP infrastructure comfortably. The platform also provides many out-of-the-box integrations with container registries, helm repos, etc., along with turnkey support for NVIDIA services such as NVIDIA NIM.

Key Rafay Platform capabilities that NCPs are leveraging to operate a PaaS experience that meets enterprise security and governance requirements are listed and defined below:

- **Kubernetes Management**

- **Virtual clusters**

- Virtual Clusters are fully working Kubernetes clusters that run on top of other Kubernetes clusters. Virtual clusters reuse worker nodes and networking of the host cluster. Each virtual cluster is deployed within a namespace and exposes its own control plane. All workloads are scheduled within the single namespace to which the virtual cluster is assigned.

- **Role-based access control (RBAC)**

- Kubernetes RBAC is a critical security control to ensure that users and workloads only have access to resources required to execute their roles. As an example, users that are allocated a namespace are automatically locked down with permissions at the namespace level only (i.e., using RoleBindings). This control ensures that users have rights only within the specific namespace and do not have the ability to perform any cluster level commands.

- **Secure remote access**

- To ensure the highest levels of security, all end users should be required to centrally authenticate using the configured Identity Provider (IdP). Once successfully authenticated, an ephemeral service account for the user is federated onto the remote cluster in a Just-in-Time (JIT) manner.

- **Network policy configuration**

- Network Policies are a mechanism to control network traffic flow within and from/to Kubernetes clusters. A common configuration network policy is to block all resources in a given namespace from exchanging network traffic with other namespaces, and from receiving traffic from outside the cluster.

- **Cluster policy configuration**

- Cluster policies can be enforced using tools such as OPA Gatekeeper and provide the means to control what users can/cannot do on the cluster. This also ensures that the clusters are always in compliance with centralized policies, such as pod privileges, label and annotation requirements, encryption requirements, limits configurations, etc. Cluster policies are closely coordinated with network policies

to ensure there is defense in depth and completeness.

- **Resource quotas**

Resource quotas are essential for managing and controlling the allocation of resources such as GPUs, CPUs, memory, and storage across multiple applications sharing a cluster. resource quotas prevent resource exhaustion, protect against overprovisioning, control resource allocation, improve performance isolation, and result in better governance.

- **Audit Logging**

A centralized and immutable audit trail of all activity performed by the users via all supported interfaces (UI and programmatic) is critical for enterprise deployments.

- **Usage metrics**

Visibility into usage of resources by tenants on a shared cluster is key to addressing chargeback and billing requirements that are fundamental to GPU Cloud operations

- **Underlay (network-level multi-tenancy) Automation**

The Rafay Platform provides automation to program and configure bare-metal servers (using Nvidia BCM API), access networks, Infiniband channels, etc., to allocate resources to NCP customers.

- **SKU Automation and Management**

- The Rafay Platform enables customers to programmatically define SKUs that consist of GPUs, CPUs, AI applications, or some combination thereof. SKUs can programmatically be offered to select customers for consumption.

- **Environment Templates Builder for 3rd Party Apps**

- The Rafay Platform also empowers NCPs to deliver internally developed or 3rd party applications to customers through Rafay's Environment Templates Builder. Once packaged as an Environment Template, apps can be added to SKUs for delivery to customers in a self-service fashion.

- **Kubernetes Cluster Lifecycle Management**

- The Rafay Platform delivers comprehensive Kubernetes cluster management capabilities, complete with fleet support for thousands of clusters. In addition to turnkey support for the CNCF-certified upstream Kubernetes distribution, the Rafay Platform also supports management for public cloud Kubernetes services (e.g., Amazon EKS, Azure AKS, Google GKE), along with proprietary Kubernetes services such as RedHat OpenShift.

- **Customer Administration**

- The Rafay Platform provides programmatic and graphical management of customers consuming GPU resources, making it easy for NCPs to manage different classes of customers through a single pane of glass.

- **Enterprise Administration**

- The Rafay Platform provides persona-specific dashboards and configuration management portals for enterprise customers to manage their deployments on top of NCP infrastructure. NCPs may choose to either white label these portals, or build their own using Rafay Platform APIs.

- **Self-Service Portal for Developers & Data Scientists**

- The Rafay Platform provides self-service portals for developers and data scientists to consume compute and AI applications on demand. NCP may choose to either white label these portals, or build their own using Rafay Platform APIs.

- **Enterprise-grade User Management**

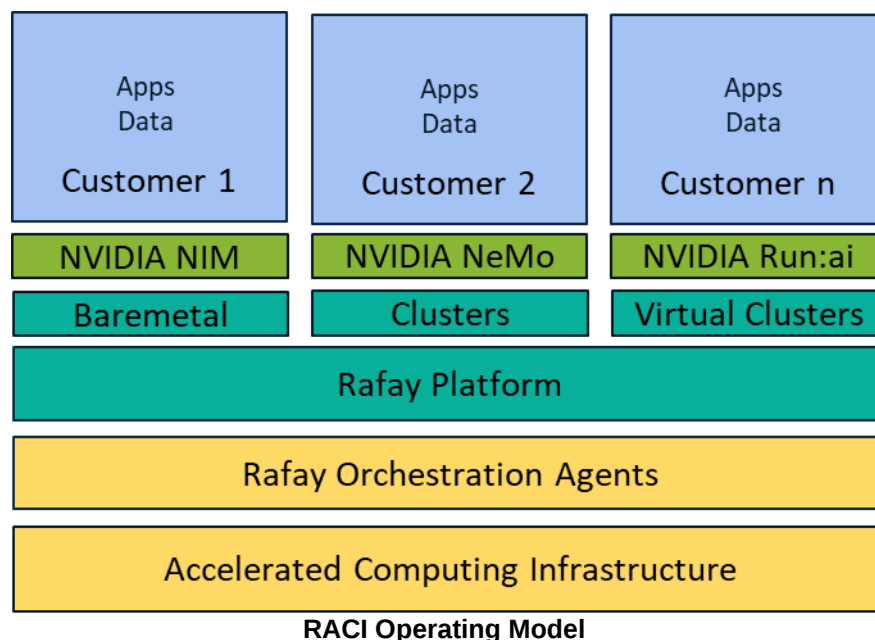
- The Rafay Platform provides extensive support for customer-specific authentication and authorization through its single sign-on (SSO) and role-based access control (RBAC) capabilities. The Rafay Platform also provides deep audit trails that can be exported to enterprise SIEMs and monitoring systems as needed.

- **Granular Usage Data for Chargeback and Billing**

- The Rafay Platform provides deep usage information across the different levels of tenants to simplify chargeback and billing workflows for NCPs. All data can be consumed via APIs or through .csv files.

Operating Model

The NCP, Rafay, and NVIDIA customers follow a shared responsibility model to manage the infrastructure, system services, application services, and customer applications and data.



- NCPs are primarily responsible for managing their accelerated computing infrastructure and running Rafay orchestration agents in their network. These orchestration agents interact with the Rafay Platform Controller to carry out cluster and resource provisioning. These components also help the NCP's platform engineering team to define and enforce compute packages and applications on the NCP hardware.
- Rafay is responsible for delivering and supporting the Rafay Platform, which consists of the Controller as well as the Rafay orchestration agents, along with the bevy of services that the Rafay Platform provides.
- NVIDIA is responsible for supporting the AI/ML and GenAI services that the NCP may use on top of the infrastructure.
- NCP customers are responsible for their applications and associated application data.

RACI

Personas

- Cloud Provider - Owner of the compute infrastructure, could be a cloud provider that is part of the NVIDIA Cloud Partner (NCP) program
- Customer - Consumer of the NCPs infrastructure
- NVIDIA - Owner of NVIDIA NIM, NVIDIA NeMo, and other services
- Rafay - Rafay Platform

RACI

- Responsible
- Accountable
- Consulted
- Inform

Roles

- Rafay Service Delivery - Implementation and Professional Services team dedicated to helping cloud provider platform engineering and SRE teams with Rafay Platform deployment, along with helping the cloud provider construct custom compute and application SKUs.
- Rafay Customer Success- Rafay's Customers Success team will interact with the cloud provider's customer support team to handle Rafay-specific issues and Level 3 support issues reported by the cloud provider's downstream customers.
- NVIDIA Customer Support - Technical support team interacting directly with customers through dedicated ticketing systems such as SFDC.
- cloud provider - Manages the underlying hardware.
- cloud provider Customer Support - cloud provider Customer Support will interact directly with the cloud provider's downstream customers for Level-1 and Level-2 support.

Activities

- Hardware Management
- Rafay Platform Management
- Rafay Orchestration Agent Management
- Kubernetes Management
- Rafay Services
- NVIDIA Services (NVIDIA NIM, NVIDIA NeMo, etc.)

Activity	Cloud Provider	Rafay Service Delivery	Cloud Provider Customer Support	Rafay Customer Support	NVIDIA Customer Support
Hardware Management					
Hardware Deployment	R				
Hardware Scale / Capacity Management	R				
Hardware Reliability	R				
Hardware Repair	R				
Ongoing Hardware Maintenance	R				
Rafay Platform Management					
Deployment	I	R	I	I	
Scale	I	R	I	I	
Service Reliability	I	R	I	I	
Bugfixes	I	R	I	I	
Ongoing Service Maintenance	C	R	I	I	
Service Security	C	R	I	I	
Rafay Orchestration Agent Management					
Agent Deployment	R	C			
Agent Reliability	R	C			

Agent Bugfix	I	R			
Ongoing Agent Maintenance	R	C			
Agent Security	I	R			
Kubernetes Management					
Kubernetes Deployment	R	C			
Kubernetes Scale / Capacity Management	R	C			
Kubernetes Reliability	R	C			
Kubernetes Bugfix	A	R			
Ongoing Kubernetes Maintenance	R	C			
Kubernetes Security	R	C			
Rafay Services					
Deployment	I	R		I	
Scale	I	R		I	
Service Reliability	I			R	
Bugfixes	I			R	
Ongoing Service Maintenance	I			R	
Service Security	I			R	
NVIDIA Services					

Deployment	R	C	I	C	C
Scale	I			R	R
Service Reliability	I				R
Bugfixes	I	I	I	C	R
Ongoing Service Maintenance	I		I	R	R
Service Security	I			I	R
SLA/SLO/SLI Definitions					
Rafay Controller SLA/SLO/SLI's		C	C	R	
Infra SLA/SLO/SLI's	R				
NVIDIA Service SLA/SLO/SLI's	I				R
Incident Response					
Escalating Nvidia Services Bugs					R
Individual Customer Communication			R		
24x7x365 Infra Support	R				
24x7x365 Service Support	R				R

Rafay Platform Deployment

The Rafay Platform can be deployed in two modes: SaaS and Self-Hosted.

- When consumed as SaaS, the management platform is hosted in Rafay's infrastructure and managed/operated by the Rafay SRE team.
- When consumed in Self-Hosted mode, the Rafay Platform is hosted in the cloud provider network in an air-gapped fashion. The Rafay Platform will be operated by the cloud provider's platform engineering and SRE team(s) with the help of Rafay's Customer Success team. The Rafay Platform can also be managed by Rafay's Service Delivery team as a fully managed service if the cloud provider prefers.

Prerequisites for SaaS Consumption

The SaaS platform is built on a zero-trust security model that only requires outbound Internet connectivity on TCP port 443 (HTTPS) from the managed clusters to the Internet-based SaaS controller for centralized management.

Network Requirements

In order to enable traffic between your Rafay agent and the Rafay Controller, you'll need access to the following URIs:

- *.rafay.dev
- *.rafay-edge.net

Security Settings

Make sure the cloud provider bare-metal servers used for deploying Kubernetes clusters have outbound connectivity on TCP port 443 to the Rafay SaaS platform hosted on *.rafay.dev.

Prerequisites for Self-Hosted Controller Deployment

Operating System

The following operating systems are supported by the Rafay Controller:

- Ubuntu 22.04 (64-bit)
- RHEL 8/9 (64-bit)

Hardware Requirements

- High Availability Controller: 3 master nodes and 1 worker node
- System Size (Minimum): 16 CPU, 64 GB RAM
- Root Disk (Minimum): 250 GB
- Temp Directory (/tmp): Minimum 50 GB (if not part of root disk)
- Data Disk (formatted): 500 GB (attached as /data volume, size may vary based on storage requirements)

Network Requirements

- For RHEL: Ensure that the nodes have connectivity to default RedHat repository servers (required for RHEL software installation and updates)
- Inbound Traffic: Allow inbound TCP traffic on port 443 to all instances and ensure all localhost ports are reachable.
- Non-DNS Environments: Enable UDP traffic on port 30,053.

Security Settings

- Configure firewall rules on your nodes to ensure the k8s and Rafay software's services can interact with each other. Please review the "Controller Deployment in Air Gapped Environments" document (Prerequisites → Infrastructure Requirements section) for additional details.

Other Considerations

- Ensure proper network connectivity between nodes.
- Verify that DNS resolution is working correctly.
- Have administrative privileges to install and configure required components.

Additional Prerequisites

Please make sure to meet the following prerequisites to install the Rafay Platform.

- NVIDIA-Certified GPU Hardware infrastructure is already installed
- The system should have a compatible OS (see list above) to install either Rafay docker agents or Rafay Kubernetes agents
- System has adequate Internet bandwidth
- DNS resolution is working properly on all the nodes

System Support

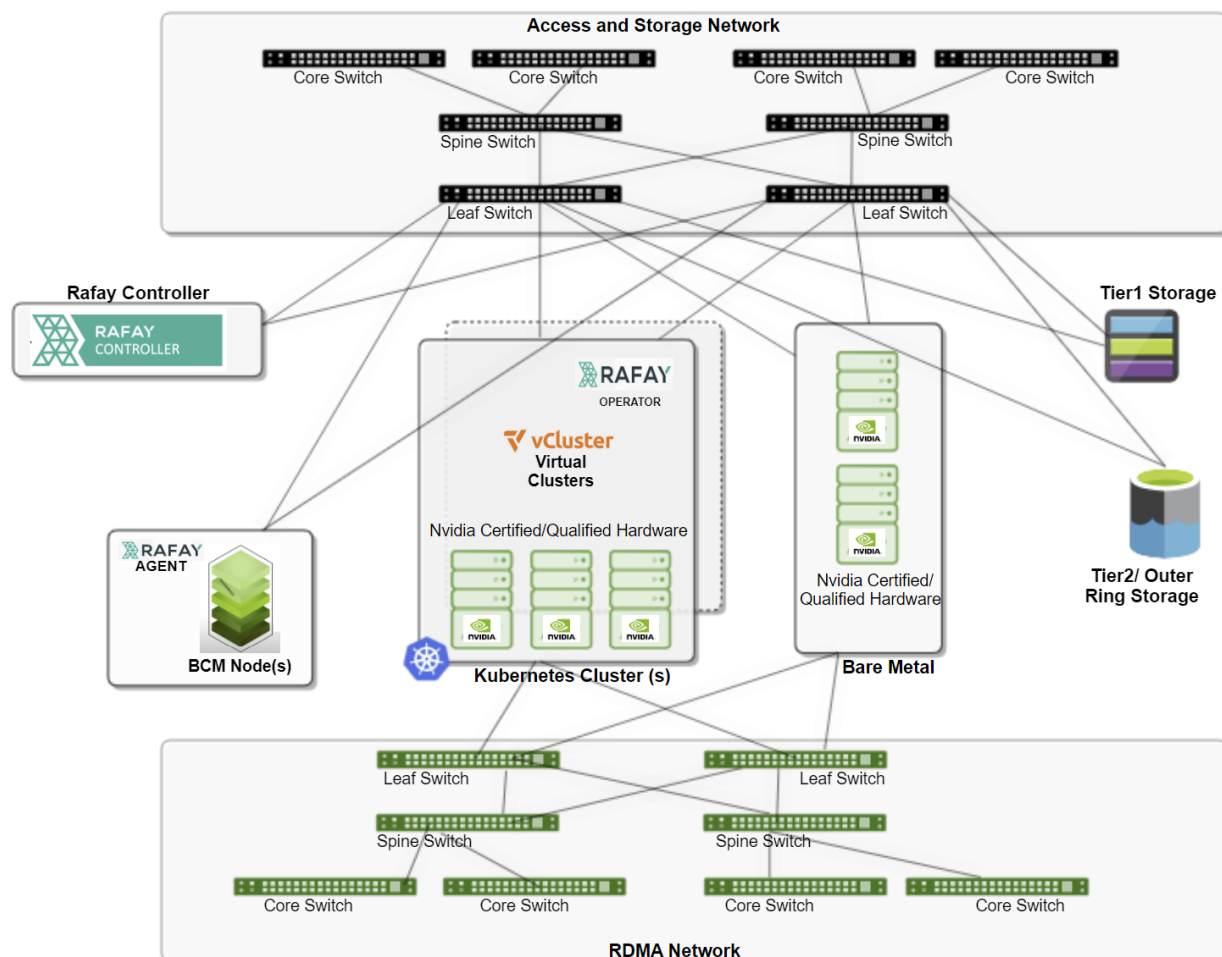
The following systems are supported by the Rafay Platform:

- You have [NVIDIA-Certified Systems](#) with NVIDIA® ConnectX® NICs for x86-64 servers
- You have [NVIDIA Qualified Systems](#) for arm64 servers (NOTE: For ARM systems, NVIDIA Network Operator is not supported)

To determine if your system qualifies as an NVIDIA-Certified System, review the list of NVIDIA-Certified Systems [here](#).

Deployment & Management

NVIDIA-Certified System Hardware and Network Topology



- Access & Storage Network:** A high-speed Ethernet network connecting users, compute, and storage, supporting low-latency access. This network is used for management tools (including the Rafay Controller, the Rafay Agents, and BCM), end-user access, and data transfers to and from storage during operations such as loading datasets or transferring model outputs. Implement this network using NVIDIA-certified/qualified Ethernet networking hardware.
- RDMA Network:** A high-throughput interconnect (InfiniBand/RoCE) for fast datapath data transfers. The RDMA network enables direct memory access between nodes without involving the operating system, reducing CPU overhead and significantly increasing data transfer speeds. RDMA is critical for workloads like distributed training, ensuring minimal latency and maximum throughput. Implement this fabric using NVIDIA-certified/qualified networking hardware.

- **NVIDIA-Certified/Qualified Compute Hardware** offered in multiple deployment options:
 - *Kubernetes*: GPU workloads can be deployed within vClusters or take over the entire host cluster. GPU resources can be allocated flexibly, either as full GPUs or partitioned (e.g. through the use of NVIDIA MiG technology).
 - *Bare Metal Nodes*: Provides direct GPU access for performance-sensitive workloads that require maximum efficiency.
- **Storage:**
 - *Tier 1*: High-performance storage for data requiring rapid access. This tier is essential for workloads like neural network training or real-time inferencing, where fast I/O and low-latency data access are critical. Actively used data, such as training datasets and model weights, reside here. Use a storage solution from the NVIDIA partner ecosystem.
 - *Tier 2/Outer Ring*: This storage serves as a repository for colder, less frequently accessed data, commonly used for archival, ETL workflows, or storing raw, unprocessed datasets. While not as fast as Tier 1, it is highly scalable and cost-efficient. Use a storage solution from the NVIDIA partner ecosystem.
- **Rafay Controller and Agent(s)**: Responsible for the lifecycle management of Kubernetes clusters, automating the provisioning and scaling of GPU-enabled clusters. It also orchestrates the deployment of workloads, facilitates policies for scheduling, monitoring, and enforcing resource quotas. It also ensures efficient GPU allocation and minimizes resource contention. Rafay's GPU slicing capabilities, leverages MiG or vGPU, ensuring optimal GPU utilization while maintaining tenant isolation.
- **BCM**: NVIDIA's Base Command Manager (BCM) manages the underlying bare metal infrastructure. BCM provides tools for provisioning and managing physical bare metal GPU nodes, handling tasks like node discovery, configuration, and lifecycle management.

Rafay Platform Configuration

The Rafay Orchestration Agent(s), in conjunction with the Rafay Platform Controller, handle(s) the lifecycle management of the compute infrastructure being provisioned and maintained through the Rafay platform. These orchestration agents are packaged as containers and can be deployed in one of two configurations: (a) within a Kubernetes cluster, or (b) as a standalone Docker container.

Please refer to the latest Rafay [documentation](#) for detailed instructions on installing Rafay Agent(s) within the cloud provider network.

Cloud Provider SKU Management

Cloud providers can offer various SKUs that reflect GPU resources and AI apps that make sense for their target customers and regional market. cloud providers can define SKUs by leveraging the template catalog supported by the Rafay Platform. To speed up delivery, cloud providers can select the SKU templates from the catalog that are closest to their preferred use cases and configure them for their needs.

Cloud Providers can also develop custom SKUs and templates using the Rafay Platform's "environment templates" framework. The following are examples of SKU templates available in the reference architecture:

- Kubernetes clusters
- Virtual Kubernetes clusters
- Baremetal servers
- Etc.

The following code snippets illustrate how cloud providers can easily deliver SKUs to their downstream customers. Note that templates are declarative and YAML based, making them human readable. Templates are expected to be stored in git repositories for "single source of truth" and version control purposes.

Kubernetes Cluster Reference SKU Template

```
apiVersion: eaas.envmgt.io/v1
kind: EnvironmentTemplate
metadata:
  name: rafay-kubernetes-cluster
  description: Manage Kubernetes cluster in Datacenter
  project: eaas
  displayName: Datacenter Kubernetes cluster
spec:
  version: v4
  resources:
    - type: dynamic
      kind: resourcetemplate
      name: oci-vm
      resourceOptions:
        version: v1
    - type: dynamic
      kind: resourcetemplate
      name: upstream-k8s-cluster
      resourceOptions:
        version: v5
      dependsOn:
        - name: oci-vm
  variables:
    - name: cluster_name
      valueType: expression
      value: $(environment.name)$
      options:
        description: Cluster Name
        override:
          type: allowed
    - name: k8s_version
      valueType: text
      value: 1.29.8
      options:
        description: k8s version
        required: true
        override:
          type: restricted
          restrictedValues:
            - 1.30.4
            - 1.29.8
  hooks: {}
  agents:
    - name: em-agents
  sharing:
    enabled: true
    projects:
      - name: demo-1
  versionState: active
  iconURL: >- https://upload.wikimedia.org/wikipedia/labs/thumb/b/ba/Kubernetes-icon-color.svg/1024px-Kubernetes-icon-color.svg.png?20210818121315
```

Virtual Cluster Reference SKU

```
apiVersion: eaas.envmgmt.io/v1
kind: EnvironmentTemplate
metadata:
  name: workspace-gpu
  labels:
    paas.envmgmt.io/objectType: computeinstances
  annotations:
    paas.envmgmt.io/cost: $1.0/GPU/hr
    paas.envmgmt.io/gpu-options: '[1,2,3]'
    paas.envmgmt.io/input-labels: '{" namespace_quota_size":"Size","gpu_quotas":"GPUs"}'
    paas.envmgmt.io/instance-sizes: >-
      [{"name": "small", "cpu":"3000m", "memory": "33Gi",
        "label":"Small"}, {"name": "medium", "cpu":"6000m", "memory": "66Gi",
        "label":"medium"}]
    paas.envmgmt.io/output-group: vcluster-import.group
    paas.envmgmt.io/output-kubeconfig: get-kubeconfig.kubeconfig
project: platform
displayName: GPU Slices
spec:
  version: v2
  resources:
    - type: dynamic
      kind: resourcetemplate
      name: vcluster-import
      resourceOptions:
        version: v1
    - type: dynamic
      kind: resourcetemplate
      name: vcluster-deploy-tshirt-sizing
      resourceOptions:
        version: v1
      dependsOn:
        - name: vcluster-import
    - type: dynamic
      kind: resourcetemplate
      name: get-kubeconfig
      resourceOptions:
        version: v1
      dependsOn:
        - name: vcluster-deploy-tshirt-sizing
  sharing:
    enabled: true
    projects:
      - name: '*'
versionState: draft
```

NVIDIA Services Deployment Configuration

With the Rafay Platform, cloud providers can provide their downstream customers with one-click deployments of NVIDIA software and application services using the service templates available in the Rafay catalog. The Rafay Platform currently supports 1-click deployments of the following NVIDIA service deployments, including:

NVIDIA NIM Deployment Template

NVIDIA NIM, part of NVIDIA AI Enterprise, is a set of easy-to-use microservices designed for secure, reliable deployment of high-performance AI model inferencing across workstations, data centers, and the cloud. Supporting a wide range of AI models, including open-source community and NVIDIA AI Foundation models, NVIDIA NIM ensures seamless, scalable AI inferencing, on-premises or in the cloud, leveraging industry-standard APIs.

NVIDIA Cloud Functions (NVCF) Deployment Template

NVIDIA Cloud Functions (NVCF) is a serverless API to deploy & manage AI workloads on GPUs, which provides security, scale, and reliability to your workloads. The API to access the workloads supports HTTP polling, HTTP streaming & gRPC.

Run:AI Deployment Template

Increase cost efficiency of Notebook Farms and Inference environments with Fractional GPUs and a custom scheduler.

NOTE: Rafay will provide additional NVIDIA services as a pre-packaged template as new NVIDIA services are made available to the market, and as per market need

NVIDIA NIM Deployment Reference Template

```
apiVersion: eaas.envmgt.io/v1
kind: environmenttemplate
metadata:
  name: nim-service
  project: platform
spec:
  version: v1
  resources:
    - type: dynamic
      kind: resourcetemplate
      name: nim-service
      resourceOptions:
        version: v1
  variables:
    - name: cluster
      valueType: text
      options:
        required: true
        override:
          type: allowed
  agents:
    - name: rafay
  sharing:
    enabled: false
  versionState: active
```

NVCF Deployment Reference Template

```
apiVersion: eaas.envmgmt.io/v1
kind: EnvironmentTemplate
metadata:
  name: nvcf-cluster-onboarding
  description: >-
    Registers Cluster with NVCF Control plane and Installs NVCF cluster agent on the cluster
  project: platform
  displayName: NVCF Cluster Agent Onboarding
spec:
  version: v3
  resources:
    - type: dynamic
      kind: resourcetemplate
      name: nvcf-onboarding
      resourceOptions:
        version: v1
    - type: dynamic
      kind: resourcetemplate
      name: nvcf-cluster-install
      resourceOptions:
        version: v1
      dependsOn:
        - name: nvcf-onboarding
  variables:
    - name: cluster_name
      valueType: text
      options:
        required: true
        override:
          type: allowed
  agents:
    - name: rafay
versionState: draft
```

NVIDIA Run:ai Deployment Reference Template

```
apiVersion: eaas.envmgmt.io/v1
kind: EnvironmentTemplate
metadata:
  name: runai-cluster-install
  project: platform
spec:
  version: v2
  resources:
    - type: dynamic
      kind: resourcetemplate
      name: runai-onboarding
      resourceOptions:
        version: v1
    - type: dynamic
      kind: resourcetemplate
      name: runai-cluster-install
      resourceOptions:
        version: v1
      dependsOn:
        - name: runai-onboarding
  variables:
    - name: cluster
      valueType: text
      options:
        required: true
        override:
          type: allowed
    - name: runai_cluster
      valueType: text
      options:
        required: true
        override:
          type: allowed
  sharing:
    enabled: false
    versionState: active
```

AI Applications Developed by Cloud Providers

Cloud providers can leverage the Rafay Platform service template framework to create and deliver AI applications that they may have developed internally. Cloud providers can also offer AI applications from their ISV partners and provide a marketplace using the Rafay Platform's catalog and marketplace integration capabilities.

Troubleshooting and Support

Troubleshooting

The Rafay Platform provides integrated diagnostics and troubleshooting capabilities for common issues. Rafay also provides users with a step-by-step guide for troubleshooting common questions when consuming the Rafay Platform, including:

- Cloud Provider customer onboarding related questions
- User management and RBAC related questions
- Cluster management related questions
- Instance management related questions (Baremetal, VM, Vcluster)
- Zero-trust access related questions
- NVCF cluster agent installation related questions
- NIM operator deployment related concerns
- NVIDIA NeMo service deployment related questions

RAFAY Support

Rafay Systems provides the following support channels:

- Slack
- Email
- Phone
- Zendesk tickets
- Technical Account Manager

The Rafay Platform provides robust documentation, all of which is available online:

- <https://docs.rafay.co>

Example Use Cases

Cloud providers can customize and deliver a wide spectrum of use cases to their customers using NVIDIA hardware and software solutions, along with the Rafay Platform. To help cloud providers go live immediately, the Rafay Platform provides a wide range of pre-built, ready-to-use solutions in the default template catalog, each of which can be packaged as SKUs. Cloud providers can additionally deliver specialized use cases by building custom templates by using Rafay's "environment templates" framework.

Following are a few example use cases:

- GPU Bare Metal Server as Service
- GPU Kubernetes Cluster as Service
- GPU Virtual Cluster as a Service
- AI Workload Deployment & Mgmt as Service
- Inference as a Service (NIM)

Supporting Documentation

- [NVCF Documentation](#)
- [NVIDIA NIM Documentation](#)
- [NVIDIA NeMo Documentation](#)
- [Rafay Docs](#)

Online documentation for the Rafay Platform.

- [Rafay Platform - GPU PaaS Deployment Guide](#)

Describes how cloud providers can deploy and operationalize the Rafay Platform on their infrastructure to deliver a GPU PaaS offering.

- [Rafay Platform - Controller Deployment in Air Gapped Environments](#)

Describes how cloud providers can deploy and configure a self-hosted version of the Rafay Platform in their infrastructure for optimal security and control.

- [Rafay Platform - Environment Manager Developer Guide](#)

Describes how cloud providers can develop, validate and use new environment templates to streamline and automate the provisioning and management of software infrastructure.

- [Rafay Platform - Multi-Tenancy Controls](#)

Describes how cloud providers can operate multi-tenant infrastructure in a secure, isolated manner, enabling them to deliver services to many more customers at lower hourly rates.