

DUMPSBOSS.

**Nutanix Certified Professional Artificial
Intelligence 6.10**

Nutanix NCP-AI

Version Demo

Total Demo Questions: 8

Total Premium Questions: 85

Buy Premium PDF

<https://dumpsboss.co>

support@dumpsboss.co

**support@dumpsboss.co
dumpsboss.co**

Topic Break Down

Topic	No. of Questions
Topic 1, AI Fundamentals	8
Topic 2, Nutanix Enterprise AI Installation and Configuration	24
Topic 3, Nutanix Enterprise AI Management	29
Topic 4, Nutanix Enterprise AI Troubleshooting	24
Total	85

QUESTION NO: 1

An administrator needs to download the Access token for installing Nutanix Enterprise AI.1

Which method should the administrator take to satisfy this task?

- A. Download from Docker hub account.
- B. Download the token from Prism Central marketplace.
- C. Download from Nutanix portal.
- D. Download from the NAI helm chart bundle.

ANSWER: C

Explanation:

Download from Nutanix portal. is correct because the access token required to install Nutanix Enterprise AI is obtained through the Nutanix customer portal. The portal is the authenticated Nutanix site used by entitled customers to access product downloads, installation artifacts, licensing-related resources, and product-specific credentials or tokens. An administrator should sign in with the organization's authorized Nutanix account, locate the Enterprise AI download or deployment resources associated with the applicable entitlement, and download or generate the required access token before beginning installation. This token allows the deployment process to authenticate to the Nutanix-hosted resources that provide the Enterprise AI software and related container artifacts. Access is therefore governed by the customer's Nutanix account and active entitlement, helping ensure that only authorized environments retrieve and deploy the product. The Nutanix Portal is available at portal.nutanix.com. Nutanix also describes its support and customer access services at [Nutanix Product Support](#), where customers can obtain authenticated access to entitled software and documentation resources.

QUESTION NO: 2

Which Nutanix AI deployment platform is supported?

- A. Native Kubernetes
- B. AWS EC2
- C. Nutanix NKE
- D. On-premises NCI

ANSWER: C

Explanation:

Nutanix NKE is a supported deployment platform for Nutanix Enterprise AI. Nutanix Enterprise AI requires a supported Kubernetes environment to host and orchestrate its containerized services, including AI model endpoints and the components used to schedule and manage GPU-enabled workloads. Nutanix Kubernetes Engine provides Nutanix-managed, CNCF-conformant Kubernetes clusters that integrate with Nutanix Cloud Infrastructure, making it an appropriate on-premises Kubernetes foundation for Enterprise AI deployments. This integration provides the Kubernetes control plane and worker-node capabilities that Enterprise AI relies on, while NCI supplies the underlying compute, storage, and networking resources. In practical deployments, administrators prepare an NKE cluster that meets the applicable Enterprise AI version's compatibility, node-sizing, storage-class, networking, and GPU requirements, then deploy the Enterprise AI services onto that cluster. Nutanix documents the supported environments and prerequisites by Enterprise AI release, so the exact supported NKE and Kubernetes versions should always be verified against the relevant compatibility documentation before installation. See the [Nutanix Enterprise AI supported environments documentation](#) and the [Nutanix Kubernetes Engine product information](#).

QUESTION NO: 3

An administrator is configuring a container registry to store custom models for secure Nutanix Enterprise AI deployment. During testing, the administrator notices the registry is accessible both via HTTP and HTTPS, with a self-signed certificate.

Which two actions should the administrator take to meet the deployment prerequisites? (Choose two.)

- A. Disable HTTP on the container registry.
- B. Install Nutanix Enterprise AI in offline mode.
- C. Replace the certificate with a trusted TLS/SSL certificate.
- D. Ensure registry runs on the same node as the cluster.

ANSWER: A C

Explanation:

"Disable HTTP on the container registry." and "Replace the certificate with a trusted TLS/SSL certificate." are required to provide the secure registry access expected for a Nutanix Enterprise AI deployment. The registry is used by the platform and its Kubernetes-based services to retrieve container images and custom model artifacts. Restricting the registry to HTTPS prevents credentials, image metadata, and model-related traffic from being transmitted through an unencrypted endpoint.

The registry certificate must also be trusted by the systems that connect to it. A certificate issued by a trusted public or enterprise certificate authority, with the appropriate registry hostname in its subject alternative names, lets cluster components validate the registry identity during TLS negotiation. If an organization uses an internal CA, its CA chain must be installed in the relevant trust stores; merely leaving an untrusted self-signed endpoint in place does not meet the normal secure-access prerequisite. This avoids image-pull and model-access failures caused by TLS verification errors while preserving authenticated, encrypted communications with the registry.

See Nutanix Enterprise AI documentation in the [Nutanix Enterprise AI Administration Guide](#) and Kubernetes guidance for [container image registries](#).

QUESTION NO: 4

Immediately after installing Nutanix Enterprise AI, an administrator runs `kubectl get pods -n nai-system` to validate the platform services.

Which two conditions indicate that a pod has completed its normal startup successfully? (Choose two.)

- A. The pod STATUS is Running.
- B. The READY value shows all containers ready.
- C. The pod STATUS is Pending.
- D. The READY value shows 1/2.
- E. The pod has an external API key assigned.

ANSWER: A B

Explanation:

A healthy pod should show a `Running` status and a `READY` value in which the number of ready containers equals the total number of containers, such as `1/1` or `2/2`. A pod in `Pending` status has not completed scheduling or initialization, while a partial `READY` value indicates that one or more required containers are not ready.

QUESTION NO: 5

Which user type is the first AI/ML admin user who installs Nutanix Enterprise AI on the Kubernetes cluster?

- A. Super User
- B. AI/ML User
- C. AI/ML Admin
- D. Super Admin

ANSWER: D

Explanation:

Super Admin is the initial privileged user type established when Nutanix Enterprise AI is installed on a Kubernetes cluster. This role bootstraps access control for the Enterprise AI environment and has the broadest administrative scope. After deployment, the Super Admin can manage the platform's user administration, including creating and managing AI/ML Admin and AI/ML User accounts, so that operational administration and regular platform access can be delegated appropriately.

The Super Admin role also provides the elevated ownership and governance capabilities needed during initial configuration. For example, it can manage shared platform resources and has authority over third-party access tokens across users, which is important for centrally administering integrations used by AI workloads. This initial role therefore establishes the trusted administrative boundary for the Enterprise AI deployment before additional users and delegated administrators are added.

Nutanix positions Enterprise AI as a Kubernetes-based platform for securely deploying and managing generative AI services and models. The Super Admin is the role that initializes the corresponding administrative model on the cluster. See the [Nutanix Enterprise AI overview](#) and the [Nutanix documentation portal](#) for product and administration documentation.

QUESTION NO: 6

Which command can the administrator use to list which versions of nai-core are available to install?

- A. `helm get repo ntnx-charts/nai-core --versions`
- B. `helm search repo ntnx-charts/nai-core --versions`
- C. `helm pull repo ntnx-charts/nai-core --versions`
- D. `helm push repo ntnx-charts/nai-core --versions`

ANSWER: B

Explanation:

The command `helm search repo ntnx-charts/nai-core --versions` lists the versions of the `nai-core` chart that are available from the locally configured Helm repository named `ntnx-charts`. Helm uses `search repo` to query the metadata cached for repositories that have been added to the local Helm client. Supplying the chart reference `ntnx-charts/nai-core` narrows the search to the Nutanix AI core chart, while the `--versions` flag instructs Helm to return all chart versions rather than only the latest matching result. This enables an administrator to identify the intended package version before performing an installation or upgrade. The repository index must be current for the displayed list to reflect newly published chart releases, so administrators commonly run `helm repo update` after adding the Nutanix chart repository or when refreshing available package metadata. Helm documents that `helm search repo` searches configured repositories and that `--versions` shows the chart versions available in the repository cache. See the official Helm command reference for [helm search repo](#) and the [helm repo update](#) command.

QUESTION NO: 7

After an endpoint is created, its serving pods remain in Pending status. Pod events show FailedScheduling messages stating that insufficient GPU resources are available.

Which two checks should the administrator perform to confirm the scheduling constraint? (Choose two.)

- A. Use `kubectl describe pod` for a pending serving pod.
- B. Review allocatable GPU capacity on the eligible worker nodes.
- C. Deactivate the API key associated with the endpoint.
- D. Change the endpoint to use a different OpenAI-compatible URL path.
- E. Delete the imported model and create a new API key.

ANSWER: A B

Explanation:

`kubectl describe pod` exposes scheduling events and the resource requests that Kubernetes could not satisfy. Reviewing the relevant worker nodes confirms whether allocatable GPU resources are exhausted or whether no compatible GPU node is available. Checking API key activity or changing the endpoint URL does not address a Kubernetes scheduler failure.

QUESTION NO: 8

An endpoint has increasing inference latency during peak demand. Its Infrastructure Summary shows sustained high CPU and memory utilization, while the cluster has sufficient unused capacity to support additional replicas.

Which two actions should the administrator take to improve endpoint performance? (Choose two.)

- A. Increase the number of endpoint instances.
- B. Allocate sufficient CPU and memory resources to the endpoint instances.
- C. Increase the number of API keys associated with the endpoint.
- D. Restart all Kubernetes worker nodes during peak demand.
- E. Disable audit events for the endpoint.

ANSWER: A B

Explanation:

The administrator should increase the number of endpoint instances so that inference requests can be distributed across additional serving replicas. The administrator should also allocate sufficient CPU and memory to each instance, because the reported utilization indicates that the current serving configuration is resource constrained. Increasing API keys does not add processing capacity, and restarting the cluster would not resolve sustained demand.