

NVIDIA-Certified Professional AI Networking

NVIDIA NCP-AIN

Version Demo

Total Demo Questions: 16

Total Premium Questions: 168

Buy Premium PDF

<https://passqueen.com>

support@passqueen.com

QUESTION NO: 1

You're troubleshooting a Spectrum-X network and notice that the System Status LED on a switch is blinking for more than 5 minutes. What is the most likely cause of this issue?

- A. The power supply unit is failing
- B. The switch is overheating
- C. The Onyx software did not boot properly

ANSWER: C

Explanation:

The Onyx software did not boot properly is correct. During a normal power-on sequence, the switch initializes its hardware and loads the network operating system. The System Status LED may blink while this initialization is in progress. However, continued blinking beyond the expected initialization interval—such as more than five minutes—indicates that the switch has not completed the transition to a running Onyx software environment. This condition is commonly associated with an interrupted or failed software load, invalid boot image, corrupted system filesystem, unsuccessful software upgrade, or a problem accessing the boot storage.

The appropriate next step is to connect through the serial console and inspect the boot output to identify where loading stops. If the switch cannot load its installed image, recovery can involve selecting a known-good boot image or using ONIE-based recovery and reinstalling the supported network operating system. NVIDIA's switch documentation describes LED indications and system-management behavior, while its networking documentation portal provides the relevant software and recovery guides for supported Spectrum switch platforms. See the [NVIDIA Networking Documentation portal](#) and the [NVIDIA Onyx documentation](#).

QUESTION NO: 2

What are two methods for accessing the operating system on a BlueField DPU?

Pick the 2 correct responses below

- A. Via the networking interfaces (data ports) in NIC mode
- B. Via the rshim interface over the PCIe bus
- C. Via the Redfish API
- D. Via rshim over a USB connection on the host

ANSWER: B D

Explanation:

Access to the BlueField DPU operating system can be established through the RShim (Remote Shim) host-side interface, which provides a management communication path between the host and the DPU. **Via the rshim interface over the PCIe bus** is the standard in-band approach when the DPU is installed in a server: the host's RShim driver exposes interfaces used to reach the BlueField management environment, including its console and network path. **Via rshim over a USB connection on the host** is also supported. USB RShim provides the same fundamental host-to-DPU management capability and is especially useful during provisioning, initial setup, or troubleshooting when PCIe-based host access is unavailable or unsuitable.

These are transport methods for the RShim management interface rather than ordinary data-plane connectivity. NVIDIA documents RShim as the mechanism that enables host communication with BlueField over PCIe or USB, and the NVIDIA-maintained RShim project describes the interfaces exposed to access the DPU. See the [NVIDIA BlueField DPU OS host-side software interfaces documentation](#) and the [NVIDIA RShim user-space repository](#).

QUESTION NO: 3

Your organization is planning to utilize Ethernet for an upcoming AI project. Spectrum-X is the selected platform for this deployment, and Adaptive Routing is a key feature.

What are the requirements included in the Spectrum-X RA for adaptive routing?

- A. SN4700, BlueField-3 SuperNIC, DDR, RoCE traffic
- B. SN5600, BlueField-3 SuperNIC, DDR, RoCE traffic
- C. SN5600, BlueField-3 SuperNIC, DDR, TCP traffic

ANSWER: B

Explanation:

SN5600, BlueField-3 SuperNIC, DDR, RoCE traffic is correct because these are the relevant platform and traffic components for Spectrum-X adaptive routing in an AI Ethernet fabric. The SN5600 is a Spectrum-4 switch platform that provides the switch-side capabilities used by Spectrum-X for high-bandwidth AI networking. BlueField-3 SuperNICs provide the host-side networking and congestion-control capabilities needed to operate efficiently with the fabric. Adaptive routing is designed to dynamically select paths according to network conditions, helping distribute AI east-west traffic across available fabric capacity rather than persistently concentrating flows on congested paths.

DDR, or dynamic distributed routing, is the routing mechanism associated with this Spectrum-X design. It works with RoCE traffic, which is central to GPU-to-GPU and storage communications in Ethernet-based AI clusters. Combining dynamic distributed routing with RoCE-aware congestion-control mechanisms enables the fabric to maintain low latency and high effective throughput for the bursty, collective communication patterns common in AI training and inference workloads. NVIDIA describes Spectrum-X as an Ethernet networking platform purpose-built for AI and highlights its adaptive routing and congestion-control capabilities. See [NVIDIA Spectrum-X AI Networking](#) and [NVIDIA BlueField-3 Documentation](#).

QUESTION NO: 4

[Spectrum-X Troubleshooting]

You're troubleshooting a Spectrum-X network and notice that the System Status LED on a switch is blinking for more than 5 minutes. What is the most likely cause of this issue?

- A. The power supply unit is failing
- B. The switch is overheating
- C. The Onyx software did not boot properly

ANSWER: C

Explanation:

The Onyx software did not boot properly is the most likely cause. On NVIDIA Spectrum switch platforms, the system-status indicator is used to communicate the switch's overall operating state, including its initialization and boot progress. A blinking system-status LED is expected while the platform performs its startup sequence, initializes hardware, and loads the network operating system. If that state persists beyond the normal startup interval, it indicates that the switch has not reached a fully operational software state. A failed or interrupted image load, filesystem problem, unsuccessful software update, or another boot-process fault can prevent Onyx from completing initialization and leave the LED blinking.

The appropriate next step is to connect through the serial console and inspect the boot output for image-loading, filesystem, or initialization errors. If the installed image cannot boot, use the supported recovery or ONIE-based installation procedure to restore a known-good NVIDIA Onyx image, then verify that the switch completes startup and reaches its normal steady system-status indication. NVIDIA's Onyx documentation provides the platform software and recovery context, while the switch hardware documentation describes status LEDs and their operational role. See [NVIDIA Onyx Documentation](#) and [NVIDIA Onyx User Manual](#).

QUESTION NO: 5

[InfiniBand Configuration]

You are validating Quality of Service behavior in an InfiniBand fabric. Which two statements about Service Levels and Virtual Lanes are true?

Pick the 2 correct responses below.

- A. A Service Level is used to classify InfiniBand traffic.
- B. The Subnet Manager maps Service Levels to Virtual Lanes.
- C. A Virtual Lane is an IP subnet identifier.
- D. Service Levels are configured only on GPU devices.

ANSWER: A B

Explanation:

A Service Level is used to classify InfiniBand traffic and the Subnet Manager maps Service Levels to Virtual Lanes are correct. The Service Level is a QoS value carried in InfiniBand communication information and is used to select the traffic treatment required by the fabric.

The fabric QoS configuration includes an SL-to-VL mapping, which determines the Virtual Lane used for traffic assigned to a particular Service Level. Virtual Lanes provide logically separate flow-control resources on a physical link and can be used to separate traffic classes and support deadlock-avoidance designs. See the [NVIDIA MLNX OFED Quality of Service documentation](#).

QUESTION NO: 6

[AI Network Architecture “ Storage Fabric]

A leading AI research center is upgrading its infrastructure to support large language model projects. The team is debating whether to implement a dedicated storage fabric for their AI workloads.

Which of the following best explains why a dedicated storage fabric is crucial for this AI network architecture?

Pick the 2 correct responses below

- A. To enable parallel data access and improve storage performance for distributed AI workloads.
- B. To ensure data security and isolation from other network traffic.
- C. To provide high-bandwidth, low-latency data access that prevents I/O bottlenecks during AI model training.
- D. To reduce the overall cost of the storage infrastructure.

ANSWER: A C

Explanation:

A dedicated storage fabric is important for large-scale AI because training performance depends on keeping GPUs continuously supplied with data. Large language model workflows read and preprocess very large datasets, often with many compute nodes accessing shared storage at the same time. “To enable parallel data access and improve storage performance for distributed AI workloads” is therefore correct: a purpose-designed storage network can deliver concurrent access at the aggregate throughput required by distributed training jobs.

“To provide high-bandwidth, low-latency data access that prevents I/O bottlenecks during AI model training” is also correct. When the storage path cannot sustain the required bandwidth or experiences congestion, GPU workers can wait for input data rather than performing computation. A dedicated fabric separates this storage traffic from other AI cluster traffic, helping provide the predictable throughput and latency needed to avoid starving costly accelerated compute resources. NVIDIA

reference architectures treat storage connectivity as a distinct network consideration and emphasize designing network and storage performance together to support GPU-accelerated workloads. NVIDIA networking technologies, including high-speed Ethernet and InfiniBand designs, support the scale-out bandwidth characteristics needed for these data-intensive environments.

References: [NVIDIA DGX BasePOD Reference Architecture](#); [NVIDIA InfiniBand Overview Documentation](#).

QUESTION NO: 7

[InfiniBand Configuration]

What are the two general user account types in MLNX-OS? Pick the 2 correct responses below:

- A. viewer
- B. monitor
- C. admin
- D. enable

ANSWER: B C

Explanation:

MLNX-OS provides the general user account types **monitor** and **admin**. These account types support role-based separation between users who need operational visibility and users authorized to manage the switch. A monitor account is intended for non-disruptive operational access: it can inspect switch status, counters, configuration, and other information needed for monitoring and troubleshooting. This permits network operations personnel to validate the health and state of the fabric without granting configuration-changing access.

An admin account has full administrative privileges. It is used for switch management tasks such as changing configuration, managing interfaces and protocols, administering users, and saving or applying system settings. NVIDIA's MLNX-OS user-management documentation identifies *admin* and *monitor* as the available user roles when creating local user accounts. Assigning these roles appropriately implements least-privilege access while retaining the operational access needed to run and support an InfiniBand environment.

See NVIDIA's [MLNX-OS User Management documentation](#) and the [NVIDIA Networking Documentation portal](#) for MLNX-OS administration guidance.

QUESTION NO: 8

[InfiniBand Troubleshooting]

An administrator suspects that a compute node has a marginal InfiniBand link. Which two counter conditions are most useful indicators of a physical-layer or signal-integrity problem?

Pick the 2 correct responses below.

- A. Increasing symbol errors
- B. Increasing link-recovery events
- C. Increasing transmitted multicast packets
- D. Increasing IP route entries

ANSWER: A B

Explanation:

Increasing symbol errors and **increasing link-recovery events** are correct. Symbol errors indicate that the receiver is detecting invalid encoded symbols on the physical link. Persistent or growing symbol-error counts can point to problems with cables, optics, connectors, lanes, or signal integrity.

Link-recovery events indicate that the port has had to retrain or recover the link after a disruption. When they occur repeatedly, they can explain intermittent connectivity or performance problems. These counters should be correlated with port state, negotiated speed and width, and cable diagnostics by using tools such as `perfquery`, `ibqueryerrors`, and `mlxlink`. See the [NVIDIA mlxlink Utility documentation](#) and the [InfiniBand Diagnostic Utilities documentation](#).

QUESTION NO: 9

Which of the following options correctly describes the difference between UFM Telemetry, UFM Enterprise, and UFM Cyber AI?

- A.** UFM Telemetry provides real-time monitoring and analysis of network performance, UFM Enterprise focuses on network management and optimization, and UFM Cyber AI detects and mitigates network security threats.
- B.** UFM Telemetry provides real-time monitoring and analysis of network performance. UFM Enterprise detects and mitigates network security threats, and UFM Cyber AI focuses on network management and optimization.
- C.** UFM Telemetry detects and mitigates network security threats. UFM Enterprise provides real-time monitoring and analysis of network performance, and UFM Cyber AI focuses on network management and optimization.
- D.** UFM Telemetry focuses on network management and optimization, UFM Enterprise detects and mitigates network security threats, and UFM Cyber AI provides real-time monitoring and analysis of network performance.

ANSWER: A

Explanation:

UFM Telemetry provides a purpose-built telemetry and visualization capability for collecting, storing, and analyzing fabric health and performance data. It gives operations teams real-time visibility into metrics such as port counters, congestion indicators, errors, and traffic behavior, enabling performance analysis and troubleshooting across an InfiniBand fabric. UFM Enterprise is the broader fabric-management platform: it provides centralized monitoring, configuration, control, automation, event management, and optimization capabilities that help administrators operate an AI-scale network reliably. UFM Cyber AI extends this operational visibility with AI-driven cyber-security analytics that identify anomalous network behavior and potential threats, helping security and network teams investigate and respond to suspicious activity. Therefore, the statement that assigns monitoring and performance analysis to UFM Telemetry, management and optimization to UFM Enterprise, and security-threat detection and mitigation to UFM Cyber AI accurately distinguishes their primary roles. NVIDIA describes the UFM product family and its fabric-management capabilities at [NVIDIA UFM](#), and provides additional information on AI-driven fabric security at [NVIDIA UFM Cyber AI](#).

QUESTION NO: 10

[Spectrum-X Configuration]

You are configuring a RoCEv2 Ethernet fabric for distributed AI training. The network must prevent avoidable packet loss for the selected RoCE traffic class.

Which two mechanisms are commonly used as part of a lossless RoCE design?

Pick the 2 correct responses below.

- A.** Priority Flow Control (PFC)
- B.** Explicit Congestion Notification (ECN)
- C.** Spanning Tree Protocol (STP)
- D.** Port Address Translation (PAT)

ANSWER: A B

Explanation:

Priority Flow Control (PFC) and **Explicit Congestion Notification (ECN)** are correct. PFC provides priority-based pause behavior, allowing a switch to pause traffic for a selected priority when buffer pressure becomes significant. In a RoCE design, PFC is applied only to the traffic priority assigned to lossless RDMA traffic rather than indiscriminately to all Ethernet traffic.

ECN marks packets when congestion begins to develop. RoCE-capable endpoints can use congestion notification derived from these marks to reduce transmission rates before queues overflow. Using ECN-based congestion control together with correctly scoped PFC helps maintain reliable RDMA operation while reducing the likelihood of widespread pause propagation. See the [NVIDIA Cumulus Linux RoCE documentation](#).

QUESTION NO: 11

[InfiniBand Troubleshooting]

Which of the following tools in Cumulus Linux is specifically useful for detecting and differentiating microbursts from regular network congestion?

Pick the 2 correct responses below

- A. Monthly network utilization reports
- B. ASIC monitoring with millisecond-level granularity
- C. SNMP polling at 5-minute intervals
- D. What Just Happened (WJH) feature for packet drop analysis

ANSWER: B D

Explanation:

ASIC monitoring with millisecond-level granularity is appropriate because microbursts are brief traffic spikes that may fill egress queues, exhaust buffer space, and cause drops within a timescale far shorter than conventional management polling. Fine-grained ASIC counters and telemetry make it possible to observe transient changes in queue occupancy, buffer utilization, packet counters, and discard counters. Seeing these short-duration changes helps operators identify bursty behavior rather than only a prolonged utilization pattern associated with sustained congestion.

What Just Happened (WJH) feature for packet drop analysis complements high-resolution monitoring by reporting packet-drop events together with the hardware-derived reason and location of the drop. In supported Cumulus Linux configurations, this visibility can identify congestion- or buffer-related drop events as they occur, providing actionable evidence that a short queue or buffer overflow caused packet loss. Combining the time-sensitive signal from ASIC telemetry with WJH drop-reason information provides both the burst signature and the cause needed to distinguish a microburst incident from ordinary longer-lived congestion. NVIDIA documents WJH as a troubleshooting facility for monitoring and analyzing dropped packets; see the [NVIDIA WJH documentation](#) and the [Cumulus Linux monitoring and troubleshooting documentation](#).

QUESTION NO: 12

[BlueField DPU Security]

You need to ensure that a BlueField DPU starts only software images that have been cryptographically authorized by the platform owner.

Which security capability provides this protection during the boot process?

- A. Secure boot
- B. ECMP hashing

C. Link Aggregation Control Protocol

D. RDMA memory registration

ANSWER: A

Explanation:

Secure boot is correct. Secure boot establishes a chain of trust by verifying the signatures of boot components before they are executed. This helps prevent unauthorized or modified firmware and software from being loaded on the BlueField platform.

For a DPU deployed as an infrastructure control point, this protection is important because networking, storage, and security services can run independently of the host operating system. Secure boot therefore helps preserve the integrity of the DPU software environment before infrastructure services begin processing traffic. NVIDIA documents platform security capabilities, including secure boot, in the [NVIDIA BlueField security documentation](#).

QUESTION NO: 13

You are deploying a Kubernetes cluster for AI workloads using NVIDIA Spectrum-X switches. You need to automate the deployment and management of networking components in this environment.

Which NVIDIA tool is specifically designed to automate the deployment and management of networking components in a Kubernetes cluster with Spectrum-X switches??

A. Mellanox OFED

B. Container Runtime

C. Network Operator

D. GPU Operator?

ANSWER: C

Explanation:

Network Operator is the NVIDIA Kubernetes operator built specifically to automate deployment, configuration, and lifecycle management of networking software components across a cluster. It uses Kubernetes custom resources to declare the desired networking configuration, then deploys and maintains the required components consistently on the applicable nodes. This model is particularly valuable for AI infrastructure because it reduces manual host-by-host setup and makes network configuration repeatable as cluster nodes are added, replaced, or updated.

For high-performance AI networking, the operator can manage elements such as NVIDIA network drivers, RDMA shared-device plugins, SR-IOV network device plugins, Multus, and related network resources. These capabilities support the accelerated networking and RDMA-oriented configurations commonly needed with NVIDIA networking hardware and AI workloads. NVIDIA's documentation describes Network Operator as a tool for deploying and managing networking components in Kubernetes, while the Spectrum-X documentation identifies its role in providing the Ethernet AI networking platform used for scaled AI deployments. See the [NVIDIA Network Operator documentation](#) and [NVIDIA Spectrum-X overview](#).

QUESTION NO: 14

[InfiniBand Congestion Control]

Which of the following statements are true about FECN and BECN in an InfiniBand fabric?

Pick the 2 correct responses below.

A. FECN is used to indicate congestion on traffic moving toward the destination.

- B. BECN is returned toward the source to provide congestion feedback.
- C. BECN assigns a new Virtual Lane to the affected flow.
- D. FECN replaces InfiniBand credit-based flow control.

ANSWER: A B

Explanation:

FECN is marked on traffic moving toward the destination when a switch detects congestion. The mark communicates that the packet encountered a congested point along its forward path.

A destination that receives FECN-marked traffic can generate BECN packets toward the source. BECN provides the reverse-path feedback that informs the source of the congestion condition. FECN and BECN are congestion-signaling mechanisms; they do not replace the credit-based flow control used by InfiniBand or define partition membership. See the [NVIDIA InfiniBand Congestion Control documentation](#).

QUESTION NO: 15

You are investigating a performance issue in a Spectrum-X network and suspect there might be congestion problems.

Which component executes the congestion control algorithm in a Spectrum-X environment?

- A. BlueField-3 SuperNICs
- B. NVIDIA DOCA software
- C. NVIDIA NetQ
- D. Spectrum-4 switches

ANSWER: A

Explanation:

BlueField-3 SuperNICs execute the Spectrum-X congestion-control algorithm. Spectrum-X is designed as an end-to-end AI Ethernet fabric in which the network can identify congestion and react dynamically to preserve high effective throughput for distributed AI workloads. The BlueField-3 SuperNIC participates at the endpoint and applies the congestion-control behavior to traffic transmission, responding to congestion feedback by regulating sending rates. This endpoint-based control is especially important for the many concurrent, high-bandwidth flows generated by GPU training and inference workloads, where unmanaged contention can produce queue buildup, latency variation, and reduced job completion efficiency.

In the Spectrum-X architecture, this SuperNIC capability works together with fabric-level telemetry and routing behavior to provide lossless, predictable AI networking at scale. NVIDIA describes Spectrum-X as combining Spectrum-4 switching with BlueField-3 SuperNIC technology and specifically identifies intelligent congestion-control capabilities as a key part of the platform. For an architectural overview, see NVIDIA's [Spectrum-X Ethernet networking page](#). Additional platform documentation is available in the [NVIDIA Spectrum-X documentation](#).

QUESTION NO: 16

You are tasked with configuring multi-tenancy using partition key (PKey) for a high-performance storage fabric running on InfiniBand. Each tenant's GPU server is allowed to access the shared storage system but cannot communicate with another tenant's GPU server.

Which of the following partition key membership configurations would you implement to set up multi-tenancy in this environment?

- A. Assign full membership to both GPU servers and storage system.
- B. Assign limited membership to both GPU servers and storage system.

- C. Assign limited membership PKey to the shared storage system and full membership PKey to each tenant's GPU servers.
- D. Assign full membership PKey to the shared storage system and limited membership PKey to each tenant's GPU servers.

ANSWER: D

Explanation:

Assign full membership PKey to the shared storage system and limited membership PKey to each tenant's GPU servers. InfiniBand P_Key membership has a deliberate asymmetric access model that is well suited to this shared-service topology. A full member can communicate with both full and limited members of the same partition. A limited member, however, can communicate only with full members and cannot communicate directly with other limited members. Therefore, configuring the shared storage endpoint as a full member allows every tenant GPU server to reach that storage service. Configuring each tenant GPU server as a limited member prevents direct GPU-server-to-GPU-server traffic between tenants, even when they share the relevant partition key.

This approach provides isolation at the InfiniBand partition layer while retaining the required storage connectivity. The storage system acts as the explicitly permitted common service endpoint, rather than creating unrestricted peer connectivity among tenant hosts. In practice, the partition must be consistently configured in the Subnet Manager/UFM policy and on the relevant host and storage adapters, with the appropriate P_Key included in their partition tables. NVIDIA's InfiniBand documentation describes P_Key partitioning and the distinction between full and limited membership: [NVIDIA MLNX_OFED InfiniBand P_Key documentation](#). NVIDIA UFM also documents centrally managed partition-key configuration for InfiniBand fabrics: [NVIDIA UFM Partition Key Management](#).