

HCIE-Data Center Network (Written) V1.0

Huawei H12-921 V1.0

Version Demo

Total Demo Questions: 23

Total Premium Questions: 231

Buy Premium PDF

<https://passqueen.com>

support@passqueen.com

Topic Break Down

Topic	No. of Questions
Topic 1, Data Center Network Architecture	26
Topic 2, Data Center Network Technologies	142
Topic 3, Data Center Network Planning and Design	10
Topic 4, Data Center Network Deployment and O&M	51
Topic 5, Mix Questions	2
Total	231

QUESTION NO: 1

Which of the following is NOT a VXLAN proprietary concept?

- A. VRF
- B. VNI
- C. VTEP
- D. NVE

ANSWER: A

Explanation:

VRF is correct because a Virtual Routing and Forwarding instance is a general routing-virtualization mechanism, not a concept created specifically for VXLAN. A VRF maintains an independent routing table and forwarding context, enabling multiple logically isolated IP routing domains to coexist on the same physical device. This capability is widely used in MPLS Layer 3 VPN deployments, enterprise network segmentation, and multi-tenant designs, regardless of whether VXLAN is deployed. VXLAN-based EVPN networks can associate tenant Layer 3 services with VRFs, but that use is an integration of two technologies rather than evidence that VRF belongs exclusively to VXLAN. The IETF's BGP/MPLS IP VPN specification describes VRF-style per-VPN routing and forwarding contexts independently of VXLAN, demonstrating that the idea predates VXLAN and has broader applicability. VXLAN itself is specified as an overlay encapsulation mechanism for extending Layer 2 segments across an IP network; its specification introduces its own overlay identifiers and endpoint behavior. See [RFC 4364: BGP/MPLS IP Virtual Private Networks](#) and [RFC 7348: Virtual eXtensible Local Area Network \(VXLAN\)](#).

QUESTION NO: 2

An MPI collection operation is when all processes in a communication space work together to accomplish a specific data operation and transmission. Which of the following items fall into a common collection action type?

- A. MPI_Scatter
- B. MPI_AllReduce
- C. MPI_Gather
- D. P2P

ANSWER: A B C

Explanation:

MPI_Scatter, **MPI_AllReduce**, and **MPI_Gather** are standard MPI collective communication routines: they are invoked by the processes in a communicator and coordinate data movement or computation across that process group. **MPI_Scatter** distributes distinct portions of data held by a root process to participating processes. **MPI_Gather** performs the complementary collection pattern, receiving data from all participating processes and assembling it at a designated root process. **MPI_AllReduce** combines values contributed by all processes using a reduction operation such as sum, maximum, or minimum, then makes the reduced result available to every process in the communicator. These routines represent common collective operation categories: data distribution, data collection, and global reduction. Collective routines are central to parallel applications because the communicator defines the participating group and the operation establishes a coordinated communication pattern for that group. The MPI standard specifies these operations as collective procedures, including their communicator-wide participation and data semantics. See the [MPI 4.1 Standard](#) and the [Open MPI documentation for MPI_Allreduce](#).

QUESTION NO: 3

Which of the following descriptions of data center interconnection solutions is wrong?

- A.** The VLAN hand-off mode enables L2 communication between data centers
- B.** In SegmentVXLAN mode, EVPN routes are transmitted in a three-stage mode to implement L3 communication between data centers
- C.** The E2E EVPNVXLAN mode directly establishes a VXLAN tunnel between the VTEPs in two data centers
- D.** After a VXLAN is established in the local mode in the segment VXLAN mode, packets are encapsulated with the local DCILeaf when they arrive at the local DCI Leaf and are sent to the peer DCI Leaf VNI, not the VNI of the peer DCI Leaf

ANSWER: D

Explanation:

The statement beginning “After a VXLAN is established in the local mode in the segment VXLAN mode” is the incorrect description because it confuses a VXLAN Network Identifier with a tunnel endpoint. In a Segment VXLAN interconnection, the DCI leaf provides a boundary between VXLAN segments. When forwarding traffic across the inter-data-center network, the local DCI leaf performs the required VXLAN processing and sends encapsulated traffic to the remote DCI leaf’s VTEP address. A VNI identifies a VXLAN segment and is carried in the VXLAN header; it is not a reachable destination to which packets can be sent. Thus, describing forwarding as being sent to a “peer DCI Leaf VNI” is technically invalid. The relevant forwarding destination is the peer VTEP, while the VNI is used to identify the tenant or broadcast domain associated with the traffic. This separation of endpoint addressing and segment identification is fundamental to VXLAN operation and to the segmented control-plane and data-plane design used for DCI. Huawei’s VXLAN documentation describes VTEPs as the tunnel endpoints and VNIs as identifiers for VXLAN networks: [Huawei CloudEngine Documentation](#). For Huawei’s product and configuration documentation portal, see [Huawei Enterprise Support Documentation](#).

QUESTION NO: 4

Which of the following describes the error in CloudFabric’s ZTP solution?

- A.** In out-of-band networking, iMaster NCE-Fabric must manage out-of-band management switches
- B.** In the in-band networking scenario, the iMaster NCE-Fabric must manage the root device
- C.** Link comparison after the device is brought online
- D.** Supports underlay network detection

ANSWER: A

Explanation:

“In out-of-band networking, iMaster NCE-Fabric must manage out-of-band management switches” is the erroneous statement. In CloudFabric zero-touch provisioning, out-of-band management switches establish the independent management network that connects devices being onboarded, servers, and iMaster NCE-Fabric. However, these management switches are not accepted by iMaster NCE-Fabric as forwarding devices to be managed through the out-of-band management network. Their underlay connectivity must instead be prepared manually. Once that management network is available, the switches undergoing ZTP can use it to reach required services such as DHCP, DNS, and the file server, and can then be discovered and onboarded by iMaster NCE-Fabric. This design keeps the out-of-band management infrastructure separate from the fabric devices that iMaster NCE-Fabric provisions and manages. Huawei describes iMaster NCE-Fabric as the platform for automated fabric deployment and lifecycle management; its product information is available at [Huawei iMaster NCE-Fabric](#). Huawei technical documentation and product guidance can also be accessed through the [Huawei Enterprise Support portal](#).

QUESTION NO: 5

What are the correct items in the following description of lossless network architecture design in multivariate computing scenarios?

- A.** In the case of a two-level network, the Spine switch does not have an uplink interface, and all ports are used to connect to the Leaf switch
- B.** When designing a network architecture, a bottom-up approach is generally used to determine how many layers of architecture to use based on the size of the network
- C.** In two-tier networking, the number of Spine switches is equal to the total number of upstream interfaces of all Leaf switches divided by the number of interfaces per Spine switch, rounded up.
- D.** First, the total number of network access interfaces is determined according to the number of computing nodes in the cluster and the number of service interfaces of each node

ANSWER: A B C D

Explanation:

In a conventional two-tier leaf–spine fabric, Spine switches form the fabric’s aggregation layer and use their interfaces to connect down to Leaf switches; there is no additional uplink tier in that design. Capacity planning is consequently performed from the access side upward. The designer first derives the required access-port count from the number of compute nodes and the number of service-facing interfaces required per node. That access requirement determines the Leaf-switch quantity and port utilization, after which the uplink capacity required from the Leaf layer is calculated.

For a non-blocking or appropriately oversubscribed fabric, the number of Spine switches must provide enough Spine-facing ports to terminate all selected Leaf uplinks. Dividing the aggregate number of Leaf uplink interfaces by the usable port count of each Spine switch and rounding upward ensures sufficient physical ports and avoids under-sizing the fabric. This bottom-up method also helps keep ECMP paths balanced and makes it practical to scale the fabric by adding Leafs or Spines as cluster capacity grows. These principles are especially important for lossless RoCE-based computing fabrics, where predictable bandwidth and carefully planned topology capacity support distributed AI and HPC traffic. Huawei’s data-center networking documentation and RoCE guidance provide further context at [Huawei Support Documentation](#) and [Huawei Data Center Network](#).

QUESTION NO: 6

Which of the following descriptions of the basic concepts of Neutron is wrong?

- A.** Security Group implements basic access control for virtual machine traffic
- B.** External Network is a special type of network that enables virtual machines to communicate with external networks
- C.** Router is used to connect subnets within a tenant to the same network or between different networks, as well as to connect internal and external networks
- D.** The virtual machine can access the external network through the floating IP, but the outside world cannot access the virtual machine through the Floating IP

ANSWER: D

Explanation:

The statement “The virtual machine can access the external network through the floating IP, but the outside world cannot access the virtual machine through the Floating IP” is the wrong description. In OpenStack Neutron, a floating IP is a routable address from an external network that is dynamically associated with a port on an instance. Neutron configures one-to-one address translation between that floating IP and the instance’s fixed IP address. This translation supports outbound connections initiated by the instance and also permits inbound connectivity from external networks to the instance through the floating IP, subject to security-group rules, network routing, and any applicable firewall policy. Consequently, a floating IP is commonly assigned specifically to make an instance reachable from outside the tenant network, for example for SSH access or publication of an application service. The ability to reach the instance is not unconditional: ingress security-group rules must allow the relevant protocol and port. However, claiming that external users cannot access the virtual machine through its floating IP contradicts the core floating-IP function. The OpenStack Neutron administration guide describes floating IPs as addresses used to access instances from external networks, and the project’s networking guide explains their association with instance ports. See [OpenStack Neutron documentation](#) and [OpenStack networking guide](#).

QUESTION NO: 7

Which of the following items can be detected by the underlay network post-event verification function implemented on iMaster NCE-Fabric?

- A. Routing black holes
- B. Network configuration parameters
- C. Routing loops
- D. Path connectivity

ANSWER: A B C D

Explanation:

iMaster NCE-Fabric uses post-event verification to validate that an underlay change has produced the intended forwarding and control-plane state. This verification covers routing black holes by checking whether traffic can still reach required destinations after a change. It also detects routing loops by evaluating forwarding behavior for paths that should be loop-free in the fabric underlay. Path connectivity is a core verification target because the controller must confirm that relevant source-to-destination paths remain reachable and usable after deployment or modification.

Network configuration parameters are also within scope because post-event verification compares the network's resulting state with the expected design and deployment intent. Checking configuration-related parameters helps identify whether devices and links are operating with the values required for the underlay to provide stable routing and forwarding. Together, these checks provide both configuration/state validation and service-oriented connectivity validation, allowing operators to identify impacts introduced by network events and changes. Huawei positions iMaster NCE-Fabric as a data-center network automation and O&M platform with automated verification capabilities; see [Huawei iMaster NCE-Fabric](#) and [Huawei Enterprise Support: iMaster NCE-Fabric](#).

QUESTION NO: 8

A transient underlay fault has been cleared before the operations team begins analysis. Which of the following information can a network snapshot in iMaster NCE-FabricInsight help preserve for subsequent fault analysis? Select three.

- A. BGP and OSPF protocol information
- B. ARP entries
- C. Packet-loss statistics on interfaces
- D. Device startup configuration files

ANSWER: A B C

Explanation:

A network snapshot captures operational network state at a selected time for later diagnosis. Relevant information includes routing-protocol state such as BGP and OSPF, ARP entries, and interface packet-loss statistics. These data help determine whether the fault involved control-plane convergence, neighbor resolution, or interface-level packet loss.

A device startup configuration is a persistent boot configuration file and is not the operational-state information collected as part of network snapshot analysis.

QUESTION NO: 9

In a lossless RoCE network, a receiver experiences congestion only for the traffic mapped to priority 3. Which of the following is the main advantage of PFC over traditional IEEE 802.3x Ethernet flow control in this situation?

- A. It pauses only the congested priority while allowing other priorities to continue forwarding.
- B. It discards packets from the congested priority before buffers become full.
- C. It reduces the sending rate by marking IP packets with ECN.
- D. It pauses all traffic on the physical interface at the same time.

ANSWER: A

Explanation:

PFC operates on a per-priority basis. The receiver can send a pause notification for priority 3 while traffic using other priorities continues to be transmitted. This reduces the head-of-line blocking that would occur if all traffic on an interface were paused.

Traditional IEEE 802.3x flow control pauses the entire link. ECN and DCQCN are congestion-notification and rate-control mechanisms; they do not provide the per-priority pause behavior described in the question.

QUESTION NO: 10

In the computing linkage scenario of the CloudFabric solution, when the VM is wired, iMasterNCEFabric senses the port on which of the host connects to the leaf point through which of the following protocols? ?

- A. LLDP
- B. NETCONF
- C. HTTP
- D. gRPC

ANSWER: A

Explanation:

LLDP is correct because the computing-linkage workflow needs iMaster NCE-Fabric to identify the physical relationship between a server host interface and the leaf-switch access port. Link Layer Discovery Protocol operates at Layer 2 and advertises device identity, interface identity, and related management information to directly connected neighbors. With LLDP information collected from the fabric, iMaster NCE-Fabric can discover which leaf interface connects to the host and use that physical attachment information when automating VM network provisioning. This is particularly important when a VM is created or attached to a network, because the controller must associate the virtual workload's host location with the appropriate access-facing leaf port before delivering network connectivity policies consistently. LLDP is therefore the discovery protocol that supplies the required host-to-leaf adjacency data in this scenario. Huawei documentation for CloudFabric describes the controller-based automated fabric and server/VM service provisioning model, while Huawei's LLDP configuration documentation describes LLDP as the protocol used to discover directly connected neighbors and their port information. See [Huawei CloudFabric documentation](#) and [Huawei LLDP configuration documentation](#).

QUESTION NO: 11

How many switches do PFC deadlock detection need to work together?

- A. Three units are synergistic
- B. It is completed independently by a single unit
- C. Two units work together
- D. Four units are synergistic

ANSWER: B

Explanation:

It is completed independently by a single unit. Priority-based Flow Control deadlock detection is a local switch function: the device monitors whether a priority queue on one of its interfaces remains continuously paused beyond the configured detection interval. The switch can identify this abnormal persistent pause state from its own received PFC state and local timers; it does not need to exchange deadlock-detection information with neighboring switches or form a multi-switch detection group. After detection, the locally configured recovery behavior can be applied, such as temporarily suppressing the affected priority's PFC handling so that traffic can resume and the deadlock can be broken. This local design is important in lossless Ethernet fabrics because a PFC pause condition can propagate across links, while each participating device must independently observe and recover the condition on its own ports. Huawei CloudEngine documentation describes PFC as an interface-level lossless Ethernet mechanism and provides local configuration and monitoring capabilities for PFC-related behavior. See the [Huawei Enterprise Documentation portal](#) and Huawei's [CloudEngine product documentation](#) for PFC configuration and feature references.

QUESTION NO: 12

In a VXLAN EVPN fabric, a remote Leaf receives an EVPN route for a tenant subnet, but the route is not installed in the corresponding tenant VRF. The BGP EVPN session is established and the route is visible in received-route information. Which of the following is the most likely cause?

- A. The tenant VRF does not import the Route Target carried by the EVPN route.
- B. The tenant VRF uses a Route Distinguisher that differs from the advertising Leaf.
- C. The VXLAN VNI is larger than the VLAN ID used at the access port.
- D. The underlay IGP does not advertise the MAC address of the remote VM.

ANSWER: A

Explanation:

An EVPN route is installed into a tenant VRF only when its Route Target matches an import Route Target configured for that VRF. The route can be received successfully through the BGP EVPN session but remain absent from the tenant routing table when Route Target import policy does not match the route's exported Route Target.

The Route Distinguisher makes otherwise identical VPN routes unique; it does not control route import. The VNI identifies the VXLAN segment and does not replace Route Target-based VPN route filtering. Underlay reachability is already sufficient because the BGP EVPN route has been received.

QUESTION NO: 13

Which concept in the MDC orchestration model is mapped to the TGW routing table in the AWS public cloud model in the CloudFabric solution?

- A. TransitRouter
- B. TGW Router
- C. TransitVAS
- D. External gateways

ANSWER: A

Explanation:

TransitRouter is the correct mapping. In CloudFabric's Multi-DC Interconnect (MDC) orchestration model, a TransitRouter represents the logical transit-routing construct used to control and distribute reachability among connected network entities. When the public-cloud side is AWS, this logical function is implemented through an AWS Transit Gateway route table. The route table determines the forwarding relationships between Transit Gateway attachments by associating attachments with a route table and propagating or statically defining destination routes. Therefore, CloudFabric uses the TransitRouter abstraction to model and orchestrate the AWS TGW route-table function consistently alongside equivalent routing constructs in other cloud environments. This abstraction lets an operator define interconnection and route-domain intent through CloudFabric rather than configuring each cloud provider's native routing object independently. AWS documents that Transit Gateway route tables contain routes that determine where traffic from attachments is directed, which is precisely the transit-routing role represented by a TransitRouter in the MDC model. See the [AWS Transit Gateway route table documentation](#) and Huawei's [CloudFabric Data Center Networking documentation](#).

QUESTION NO: 14

Which of the following descriptions of the business chain is wrong?

- A. In the NSH implementation, the service chain function is implemented by inserting the NSH header into VXLAN packets
- B. In the PBR implementation, the controller delivers PBR diversion policy routes to SFF and SC nodes hop-by-hop through NETCONF
- C. In PBR implementation, the controller delivers filtering and redirection policies to the SC nodes to selectively redirect service traffic
- D. In the NSH implementation, the SF node must support NSH

ANSWER: D

Explanation:

"In the NSH implementation, the SF node must support NSH" is the incorrect description. Network Service Header is designed to carry service-path and service-index metadata that identifies the required service-chain traversal. Although an NSH-aware service function can process NSH information directly, NSH does not require every service function to be NSH-capable. A service-function proxy can act as the NSH-aware attachment point for a legacy or NSH-unaware service function: it removes or adds the NSH encapsulation as necessary while forwarding the payload to and from that service function. This enables existing physical or virtual appliances, such as firewalls and load balancers that do not implement NSH, to participate in an NSH-based chain without modification. Therefore, NSH support is required at the relevant service-chain forwarding/proxy functions that handle NSH encapsulation and forwarding, but it is not an absolute requirement of the SF node itself. This separation is a core interoperability benefit of the NSH architecture. See the IETF NSH architecture definition in [RFC 7665](#) and the NSH protocol specification in [RFC 8300](#).

QUESTION NO: 15

Which of the following descriptions of BD in a VXLAN network is correct?

- A. BD messages are encapsulated in VXLAN packets to identify a large Layer 2 broadcast domain
- B. In a VXLAN network, the BD and the VLAN tag of the connected service packet form a 1:N mapping
- C. In a VXLAN network, BD is a local concept, and as long as the L2 VNI is the same on different leaves, it can be considered to be in the same large Layer 2 broadcast domain
- D. Multiple BDs can share a single VBDIF interface

ANSWER: C

Explanation:

In a Huawei VXLAN deployment, a bridge domain (BD) is a local logical Layer 2 broadcast-domain construct on a device. It provides the context in which local access-side traffic is learned, bridged, and associated with VXLAN forwarding

information. The VXLAN Network Identifier (VNI), specifically the Layer 2 VNI for Layer 2 VXLAN service forwarding, is the overlay identifier carried in VXLAN encapsulation. Therefore, when corresponding BDs on different leaf switches are bound to the same Layer 2 VNI, those leaves extend the same tenant Layer 2 broadcast domain across the VXLAN overlay. Frames entering the BD can be encapsulated with that VNI and transported between VTEPs, while remote VTEPs use the same VNI association to place received traffic into their corresponding local service context. This separation is important: BDs are locally configured service instances, whereas the common L2 VNI identifies the shared overlay segment between VTEPs. Huawei's VXLAN configuration documentation describes associating a BD with a VXLAN network identifier to implement Layer 2 service extension. See Huawei's [VXLAN Configuration Guide](#) and the [Huawei CloudEngine VXLAN documentation](#).

QUESTION NO: 16

In an enterprise data center, the customer assigns a virtual firewall (VSYS1 and VSYS2) to the finance department (VPC1) and the sales department (VPC2), VSYS1 If it is located in the same ServiceLeaf as VSYS2, what are the correct descriptions of traffic orchestration in this scenario?

- A.** In the scenario where the Border Leaf and Service Leaf are separated, the traffic between the VPC and the VPC needs to pass through the Border Leaf
- B.** In the scenario where the Border Leaf and Service Leaf are separated, the controller will create VPC1 and VPC2 on the Service Leaf according to the VPC interoperability requirements Associated VRF
- C.** The tenant needs to specify the interworking VRF on the controller to complete the VPC interworking in this scenario
- D.** Regardless of whether the Border Leaf and Service Leaf are co-located or separated, the traffic model is consistent

ANSWER: B D

Explanation:

In a service-leaf-separated design, traffic requiring firewall service insertion is steered to the Service Leaf that attaches the firewall interfaces and the relevant VSYS instances. For inter-VPC communication, the controller therefore provisions the required associated VRF context on the Service Leaf for VPC1 and VPC2, allowing the service chain to direct traffic into the appropriate firewall virtual systems while maintaining tenant routing isolation. The tenant expresses the intended VPC connectivity and security-service requirements through the controller; the fabric controller translates that intent into the required underlay and service-node forwarding configuration.

The statement that the traffic model remains consistent whether the Border Leaf and Service Leaf roles are co-located or separated is also correct. Co-location changes where functions reside physically, but it does not alter the logical service-chaining behavior: traffic is classified in the tenant VRF, directed through the firewall service, and returned to the appropriate tenant forwarding context. The selection and placement of VSYS instances on the firewall/service-leaf side determines the relevant service path, rather than requiring a different tenant-facing traffic model. Huawei describes VRF-based tenant isolation and inter-VPC connectivity concepts in its [CloudFabric documentation](#) and provides Data Center Network solution guidance at [Huawei Data Center Network](#).

QUESTION NO: 17

In a CloudFabric solution, the fault detection and fault-tolerant design of the firewall includes ?().

- A.** A heartbeat cable is deployed between the two FWs to check the status of the peer FW and the heartbeat link connectivity
- B.** It is recommended that each FW and service be connected in a single-arm mode, and the upstream and downstream logical ports fail at the same time, resulting in higher failover performance
- C.** Link fault: If any link on the edge leaf of the FW primary connection fails, the M-lag member port goes down, and the double homing becomes single homing, which does not affect the forwarding of FW primary traffic
- D.** FW fault: After the FW is switched over to the primary/standby, the FW sends a free ARP packet to trigger the refresh of the forwarding entries on the switch and direct the traffic to the FW standby

ANSWER: A C D

Explanation:

A heartbeat cable deployed between the two firewalls is part of the high-availability design because it provides a dedicated channel for the devices to monitor peer status and heartbeat connectivity. This monitoring supports timely active/standby fault detection and switchover decisions. In the CloudFabric firewall attachment design, M-LAG dual homing protects firewall connectivity: when one member link on the edge-leaf side fails, the remaining member link continues to provide the forwarding path. The firewall therefore remains reachable through a single active attachment while primary service traffic continues to be forwarded. After an active/standby firewall switchover, sending a gratuitous ARP packet is also essential to rapid service recovery. The packet updates the relevant Layer 2 and Layer 3 forwarding information on connected switches, so traffic that previously reached the former active firewall is redirected to the newly active firewall without waiting for normal ARP aging. Together, peer heartbeat detection, M-LAG link resiliency, and gratuitous-ARP-based forwarding refresh provide the expected fault-detection and fault-tolerance behavior for a CloudFabric firewall service chain. Huawei describes CloudFabric as an architecture built for highly available data-center networking and provides firewall product and configuration resources through its [CloudFabric product page](#) and [Next-Generation Firewall support portal](#).

QUESTION NO: 18

Which of the following scenarios is CloudFabric suitable for lossless networking?

- A. HPC high-performance computing scenarios
- B. AI Distributed computing scenarios
- C. Distributed storage scenarios
- D. Centralized storage scenario

ANSWER: A B C D

Explanation:

CloudFabric lossless networking is designed for data-center workloads that are highly sensitive to packet loss, congestion, and latency. HPC high-performance computing scenarios are suitable because parallel compute jobs exchange large volumes of data among nodes; lossless Ethernet helps maintain efficient, predictable communication for these tightly coupled workloads. AI Distributed computing scenarios also require lossless networking, particularly for distributed training, where frequent parameter synchronization and collective communication can be severely affected by dropped packets and retransmissions. Distributed storage scenarios benefit because storage traffic between clustered nodes requires dependable, high-throughput transport to preserve performance and service continuity. Centralized storage scenario deployments are equally applicable: hosts and storage systems can use lossless Ethernet to carry storage traffic with low loss and controlled congestion. CloudFabric supports lossless Ethernet technologies and congestion-management capabilities for these workloads, including RoCE-oriented data-center designs. Huawei describes CloudFabric as a data-center network solution for cloud, computing, and storage infrastructure, with high-performance network capabilities for modern workloads. See Huawei's [CloudFabric solution overview](#) and its [data center network solutions](#) for further context on these application scenarios.

QUESTION NO: 19

Which of the following descriptions of the open source OpenStack Glance component is correct?

- A. Modifying an instance that is started based on an image is actually modifying the image
- B. Glance provides the ability to discover, register, and retrieve virtual machine images
- C. Virtual Machine Instance images provided by Glance can only be placed in object storage provided by Swift
- D. You can launch only one instance from the same image

ANSWER: B

Explanation:

Glance provides the ability to discover, register, and retrieve virtual machine images. In OpenStack, the Image service is the centralized catalog and delivery service for images used to create compute instances. It stores image metadata, lets users and services search or discover available images, supports image registration and upload, and makes image data available to Nova when an instance is launched. The image is a reusable template, typically containing an operating system and optionally preinstalled software or configuration. This separation lets cloud administrators publish controlled base images while tenants can select an appropriate image during instance provisioning. Glance supports multiple storage back ends through its store drivers, so its fundamental role is image management and image retrieval rather than being limited to a specific storage technology. It also supports image metadata and visibility controls, which help determine how images are found and shared within an OpenStack deployment. For the OpenStack Image service overview and its documented capabilities, see the [OpenStack Glance documentation](#) and the [Glance architecture guide](#).

QUESTION NO: 20

Based on the multiple computing scenarios of Huawei's CloudFabric intelligent lossless network solution, what are the correct descriptions of HPC cluster network design?

- A.** The HPC cluster network adopts a planar design that considers the business plane, management plane, and computing plane separately
- B.** High-performance computing networks are used for communication between computing nodes, requiring networks with low latency and high throughput, and are generally deployed as lossless networks
- C.** The service management network is used for cluster management, resource monitoring, and other related business operations, and is usually deployed as a TCP/IP lossy network
- D.** Both the high-performance computing network and the storage front-end network must be deployed independently to ensure high network performance and prevent mutual interference

ANSWER: A B C**Explanation:**

HPC cluster networking separates traffic according to its performance requirements and operational purpose. The planar design that considers the business plane, management plane, and computing plane separately is appropriate because it provides clear traffic boundaries and lets each plane be designed and operated according to the workloads it carries. In particular, communication among compute nodes is highly sensitive to latency, throughput, and packet loss. A high-performance computing network for this traffic is therefore designed as a high-bandwidth, low-latency lossless network, typically using lossless Ethernet mechanisms and congestion-control capabilities to maintain efficient distributed-computing performance.

The service management network has a different role: it carries cluster administration, resource monitoring, and related management services. These applications generally use ordinary TCP/IP communications and do not require lossless transport, so a conventional lossy IP network is suitable for this plane. This separation enables the HPC fabric to focus its deterministic performance resources on compute traffic while retaining a practical and manageable network for operational services. Huawei positions CloudFabric as a data-center network architecture for high-performance, intelligent networking, including lossless network capabilities; see [Huawei CloudFabric](#) and [Huawei Data Center Networking solutions](#).

QUESTION NO: 21

Which of the following options are included in the action type of a runbook in the CloudFabric network business open programmable overall solution?

- A.** Rollback
- B.** Withdraw
- C.** Commit
- D.** DryRun

ANSWER: A B C D

Explanation:

CloudFabric's network-business open-programming framework uses runbooks to describe the executable lifecycle operations of a network service. The listed action types—**Rollback**, **Withdraw**, **Commit**, and **DryRun**—are all runbook actions. **Commit** applies the intended service or configuration change, while **DryRun** supports a pre-execution validation or simulation workflow so that the change can be assessed before it is formally applied. **Rollback** is the lifecycle action used to restore the relevant prior state after a change needs to be reverted. **Withdraw** supports removal of an already delivered service or its associated network configuration through the runbook lifecycle. Together, these actions enable controlled service delivery, validation, reversal, and decommissioning in programmable data-center network operations. Huawei positions CloudFabric as an automated, intent-driven data-center network solution, and its product and documentation resources provide the broader context for this lifecycle-oriented automation model. See Huawei's [CloudFabric product page](#) and the [Huawei CloudFabric documentation portal](#).

QUESTION NO: 22

After the dual-stack transformation of the Internet access area of the data center network, which of the following methods can the network egress choose to interconnect with the operator 's IPv network?

- A. Static IPv6 routes
- B. OSPFv3
- C. IBGP4+
- D. EBGp+

ANSWER: A D

Explanation:

Static IPv6 routes and EBGp+ are appropriate choices for IPv6 interconnection between a data-center Internet egress and an operator. Static IPv6 routes provide a simple, deterministic approach when the number of external prefixes and links is small. The egress device can install a configured IPv6 default route or specific IPv6 routes toward the provider, while return routes are configured by the provider as required. This approach is commonly used where routing policy and resiliency requirements are limited.

EBGP+ is the IPv6 address-family capability of external BGP and is the standard dynamic-routing approach for an inter-domain connection to an ISP. It allows the enterprise and operator to exchange IPv6 reachability information, apply import and export routing policies, control advertised prefixes, and support scalable multihoming and failover. BGP extensions for multiple network-layer protocols define the use of an IPv6 address family in BGP; see [RFC 4760](#). Huawei documentation also identifies BGP as a protocol for exchanging routing information between autonomous systems; see the [Huawei BGP Configuration Guide](#). Accordingly, a static IPv6-routing design or an EBGp+ peering design can be selected according to the scale, redundancy, and routing-policy requirements of the operator connection.

QUESTION NO: 23

Which of the following needs to be configured in OpenStack when configuring intranet cross-VPC access (over the wall) in OpenStack? ?

- A. Networking
- B. Routing
- C. VPC interworking
- D. Specify the source firewall

ANSWER: A B

Explanation:

Networking and **Routing** are required foundations for intranet connectivity between isolated VPC environments. Networking defines the relevant virtual networks and IP subnets, as well as the interfaces through which workloads attach to those networks. Each VPC normally represents an isolated Layer 3 domain, so the participating address spaces and network interfaces must be planned and configured before traffic can be exchanged.

Routing provides the forwarding path for traffic whose destination is located in the other VPC. The route configuration must include an appropriate destination CIDR and a reachable next hop or routing device for the cross-VPC path. Return routing is equally important: both directions require valid forwarding information so that response traffic follows a routable path back to the initiating workload. This reflects the standard Neutron model in OpenStack, where networks, subnets, router interfaces, and routing behavior provide Layer 3 communication between otherwise separate tenant networks. Huawei Cloud VPC documentation likewise identifies route tables as the mechanism that controls traffic forwarding within and beyond a VPC.

See the [OpenStack Neutron networking introduction](#) and the [Huawei Cloud VPC User Guide](#) for the relevant virtual-networking and route-table concepts.