

DUMPS ARENA

VMware Cloud Professional

VMware 2v0-33.22pse

Version Demo

Total Demo Questions: 10

Total Premium Questions: 106

Buy Premium PDF

<https://dumpsarena.co>

sales@dumpsarena.co

sales@dumpsarena.co
dumpsarena.co

QUESTION NO: 1

In VMware Cloud, who is responsible for the encryption of virtual machines?

- A. Native cloud provider
- B. Customer
- C. VMware Cloud Provider Partner (VCPP)
- D. VMware

ANSWER: B

Explanation:

Customer is correct because VMware Cloud uses a shared-responsibility model: VMware operates and secures the underlying cloud service and physical infrastructure, while the customer remains responsible for workloads deployed in its software-defined data center. Virtual machines, their guest operating systems, applications, and the data they contain are customer-managed components. Consequently, when encryption is required within a virtual machine—such as guest operating-system disk encryption, application-level encryption, database encryption, or management of encryption keys and security policies—the customer must select, configure, operate, and maintain those controls.

This responsibility also includes ensuring that encryption settings align with organizational security requirements, access-control policies, key-management practices, and regulatory obligations. VMware Cloud provides the platform capabilities and protects service infrastructure, but it does not make workload-specific encryption decisions or administer guest-level security configurations on the customer's behalf. VMware's VMware Cloud on AWS documentation describes the shared-responsibility approach and identifies customer responsibility for securing and managing virtual machines and guest security. See the [VMware Cloud on AWS encryption documentation](#) and the [AWS Shared Responsibility Model](#).

QUESTION NO: 2

A cloud administrator must allow application servers to communicate with database servers on TCP port 1433 in a VMware Cloud SDDC. Both server groups are identified dynamically by NSX tags, and all other east-west traffic to the database servers must remain denied.

Which two NSX configurations should the administrator use? (Choose two.)

- A. Create NSX Groups with dynamic membership criteria based on the application and database tags.
- B. Create a distributed firewall rule allowing TCP port 1433 from the application group to the database group.
- C. Create a gateway firewall rule allowing TCP port 1433 from the application group to the database group.
- D. Create a routed segment for the database servers.
- E. Create a DNAT rule for each database server.

ANSWER: A B

Explanation:

NSX Groups with dynamic membership criteria allow the administrator to define application-server and database-server membership from tags. This ensures that security policy automatically applies when virtual machines are added, removed, or retagged.

A **distributed firewall rule** should then allow TCP port 1433 from the application-server group to the database-server group. Distributed firewall enforcement occurs at the workload virtual NIC, making it appropriate for granular east-west controls. A gateway firewall is primarily used for traffic traversing a gateway, while a routed segment or NAT rule does not enforce the required workload-to-workload security policy.

QUESTION NO: 3

A cloud administrator wants to view and manage workloads across both an on-premises environment and a VMware Cloud on AWS software-defined data center (SDDC).

Which solution meets this requirement?

- A. Enhanced Linked Mode
- B. VMware HCX
- C. vCenter Single Sign-On
- D. Hybrid Linked Mode

ANSWER: D

Explanation:

Hybrid Linked Mode is the VMware Cloud on AWS capability designed to provide unified visibility and administration across an on-premises vSphere environment and a cloud SDDC. It connects the on-premises vCenter Server environment to the vCenter Server instance in VMware Cloud on AWS, allowing an administrator to use the vSphere Client to view inventory and manage workloads in both locations through a common management experience. This is particularly valuable in hybrid-cloud operations because administrators can work with virtual machines, hosts, clusters, networks, and related inventory without treating the cloud SDDC as an entirely separate management silo. Hybrid Linked Mode uses vCenter Server Single Sign-On federation to establish the trusted relationship between the environments, while delivering the cross-environment inventory visibility and management required by the scenario. VMware also documents Hybrid Linked Mode as an integration that supports hybrid operations between on-premises vSphere and VMware Cloud on AWS. For implementation concepts and requirements, see the [VMware Cloud on AWS Hybrid Linked Mode documentation](#) and the [vCenter Server management guidance for VMware Cloud on AWS](#).

QUESTION NO: 4

Which three types of gateways can be found in VMware cloud on AWS (Choose three?)

- A. Distributed Tier-1
- B. Standard Tier-1
- C. Tier-0
- D. Compute Tier-1
- E. Management Tier-1
- F. Management Tier-0

ANSWER: C D E

Explanation:

VMware Cloud on AWS uses a gateway design that separates workload connectivity from the infrastructure-management plane while providing north-south routing. The **Compute Tier-1** gateway, commonly called the Compute Gateway (CGW), provides network services and connectivity for workload segments. It is where workload virtual machines connect to routed networking and where workload networking policies and services are applied. The **Management Tier-1** gateway, or Management Gateway (MGW), serves the SDDC management network, supporting connectivity for VMware-managed components such as vCenter Server, ESXi hosts, and NSX management services. This separation helps preserve the distinct security and operational requirements of management and customer workload traffic.

Both Tier-1 gateway domains connect upstream to the **Tier-0** gateway. The Tier-0 gateway provides the shared north-south routing layer for the SDDC and connects its Tier-1 gateways to external networks, including connections through VMware Transit Connect, Direct Connect, VPN, and internet-facing services as configured. Together, the Compute Tier-1,

Management Tier-1, and Tier-0 gateways form the core VMware Cloud on AWS logical routing architecture. VMware's networking overview and gateway documentation describe this management/compute separation and the Tier-0-to-Tier-1 topology: [VMware Cloud on AWS Networking 101](#) and [VMware Cloud on AWS documentation](#).

QUESTION NO: 5

On VMware Cloud on AWS, which type of host do you use when you require high local storage requirements and additional cores for your workloads? (Select one option)

- A. ve-standard-72
- B. i3en.metal
- C. i3.metal
- D. AV36

ANSWER: B

Explanation:

i3en.metal is the appropriate VMware Cloud on AWS host type for workloads that need both substantial local NVMe storage capacity and a higher physical-core count. This storage-optimized instance type provides 48 physical CPU cores, 768 GB of memory, and approximately 45.84 TiB of local NVMe SSD capacity. Those resources make it well suited to data-intensive workloads whose performance depends on high-throughput, low-latency local storage, such as large databases, analytics platforms, and applications with significant caching or data-processing requirements. The additional cores also provide greater compute density for workloads that need more CPU resources while retaining the large local-storage footprint. In VMware Cloud on AWS, local instance storage is consumed by vSAN, so the host's NVMe capacity directly supports the available vSAN datastore capacity for the SDDC. VMware's VMware Cloud on AWS documentation provides current operational guidance and platform information at [VMware Cloud on AWS Documentation](#). The underlying AWS instance-family specifications, including the storage-optimized characteristics of i3en instances, are available at [AWS i3en Instance Types](#).

QUESTION NO: 6

An application team must run a machine-learning workload only on Kubernetes worker nodes that have the label *workload-type=gpu*. The GPU nodes are tainted so that ordinary workloads are not scheduled on them.

Which two Pod configuration settings are required? (Choose two.)

- A. A node selector that matches the GPU node label
- B. A toleration that matches the GPU node taint
- C. A HorizontalPodAutoscaler for the workload
- D. A resource request for CPU and memory only
- E. A Service that selects the workload Pods

ANSWER: A B

Explanation:

A **node selector** directs the Kubernetes scheduler to place the Pod only on nodes with the required label. In this case, the selector matches *workload-type=gpu*, ensuring that the workload is considered only for GPU worker nodes.

A matching **toleration** is also required because the target nodes are tainted. The toleration permits this Pod to be scheduled onto nodes carrying the specified taint. A resource request can influence placement when resources are scarce, but it does not constrain placement to labelled GPU nodes or override a taint. An affinity rule could also express placement preferences or requirements, but it does not by itself tolerate the node taint.

QUESTION NO: 7

A cloud administrator is asked to evaluate a number of disaster recovery solutions for the business. The current on-premises environment is built around the latest version of VMware vSphere 7.0.

The following requirements must be met:

- Follow an on-demand cloud consumption model
- Must be a managed offering
- Deliver a recovery point objective (RPO) of no more than 30 minutes
- Rapid power-on of recovered virtual machines/ assuming cloud capacity availability
- Must accommodate for single region failure

Which solution would meet these requirements?

- A. VMware Cloud Disaster Recovery
- B. VMware Cloud on AWS Stretched Cluster
- C. VMware vSphere Replication
- D. VMware Site Recovery Manager

ANSWER: A

Explanation:

VMware Cloud Disaster Recovery is the appropriate solution because it is a VMware-managed disaster-recovery-as-a-service offering that uses an on-demand cloud consumption model. It protects compatible vSphere workloads through VMware Cloud on AWS and provides cloud-based recovery capabilities without requiring the organization to permanently operate a fully provisioned secondary recovery site. This aligns with the requirement to consume recovery resources as needed, while still enabling rapid recovery when cloud capacity is available.

The service uses high-frequency snapshots and can support recovery point objectives as low as 30 minutes. Its recovery workflow supports powering on protected virtual machines in a recovery SDDC, allowing workloads to be brought online quickly during a disaster event. VMware Cloud Disaster Recovery also includes resiliency capabilities designed for recovery scenarios involving an Availability Zone or a single-region disruption, supporting the stated regional-failure requirement. The managed service model, vSphere workload compatibility, on-demand recovery infrastructure, and 30-minute RPO capability together make VMware Cloud Disaster Recovery the solution that meets the complete set of requirements. See the VMware documentation on [availability-zone failure handling](#) and the [VMware Cloud Disaster Recovery release notes](#).

QUESTION NO: 8

A cloud administrator would like the VMware Cloud on AWS cluster to automatically scale-out and scale-in based on resource demand. Which two Elastic DRS policies can be configured to meet this requirement? (Choose two.)

- A. Elastic DRS Baseline policy
- B. Optimize for Best Performance policy
- C. Optimize for Lowest Cost policy
- D. Custom Elastic DRS policy
- E. Optimize for Rapid Scale-Out policy

ANSWER: B C

Explanation:

Optimize for Best Performance policy and **Optimize for Lowest Cost policy** are Elastic DRS policy presets that automate VMware Cloud on AWS host-capacity adjustments in response to cluster resource demand. Elastic DRS continuously evaluates the SDDC cluster's CPU, memory, and capacity requirements, then requests host additions when demand requires more capacity and removes hosts when sustained demand permits reducing capacity. This supplies elasticity without requiring the cloud administrator to manually add or remove hosts for routine demand changes.

Optimize for Best Performance uses a more performance-oriented capacity posture, maintaining additional available headroom so workloads are less likely to experience resource contention as demand grows. Optimize for Lowest Cost uses a more cost-conscious posture, allowing capacity to be used more efficiently before expanding the cluster and scaling in when excess capacity is no longer needed. Both policies therefore meet the stated requirement for automated scale-out and scale-in; the practical choice is driven by whether workload performance headroom or infrastructure-cost efficiency is the primary operational objective. Elastic DRS respects the applicable cluster sizing and operational constraints while making these capacity decisions. See VMware Cloud on AWS documentation at [VMware Cloud on AWS Technical Documentation](#) and the AWS overview at [VMware Cloud on AWS on AWS Docs](#).

QUESTION NO: 9

Which two Tanzu Kubernetes Grid service component must an administrator configure within VMware Cloud to enable to deploy a namespace or their Kubernetes Application developments? (Choose two)

- A. Tanzu Service Mesh
- B. Tanzu Application Platform
- C. Tanzu Kubernetes Cluster
- D. Management cluster
- E. Tanzu Observability by Wavefront

ANSWER: C D

Explanation:

Tanzu Kubernetes Cluster and **Management cluster** are the core Tanzu Kubernetes Grid components used to establish and operate Kubernetes environments. The management cluster provides the management plane for Tanzu Kubernetes Grid: it is responsible for lifecycle operations such as creating, scaling, upgrading, and deleting the Kubernetes clusters that run workloads. It also supplies the management services and configuration required to consistently operate those clusters.

A Tanzu Kubernetes Cluster is the workload Kubernetes cluster on which development teams deploy Kubernetes resources and applications. After the management environment is available, administrators can provision and manage these workload clusters according to the organization's requirements. The workload cluster supplies the Kubernetes control plane and worker capacity for application deployment, while the management cluster maintains the lifecycle and centralized operational control of that environment. Together, these components implement the fundamental Tanzu Kubernetes Grid model of a management cluster that manages one or more workload clusters.

VMware's Tanzu Kubernetes Grid documentation describes this architecture and the management-cluster role in cluster lifecycle management. See [VMware Tanzu Kubernetes Grid documentation](#) and [VMware Tanzu Kubernetes Grid product overview](#).

QUESTION NO: 10

How is a Tanzu Kubernetes cluster deployed in a VMware Cloud environment?

- A. Using the VMware Cloud Console
- B. Using VMware Tanzu Mission Control
- C. Using the standard open-source kubectl
- D. Using the vSphere PlugIn for kubectl

ANSWER: A

Explanation:

Using the VMware Cloud Console is correct because the console is the VMware Cloud management interface used to provision Tanzu Kubernetes clusters in the applicable VMware Cloud environment. An administrator begins in the Tanzu Kubernetes Grid area of the console, chooses the target software-defined data center, and supplies the cluster configuration. The deployment workflow captures key cluster settings, including the Kubernetes version, control-plane and worker-node sizing, networking, and access configuration. VMware Cloud then provisions the required Kubernetes infrastructure on the selected cloud SDDC and exposes the resulting cluster for ongoing Kubernetes administration.

This console-driven process provides an integrated cloud workflow: it connects cluster provisioning to the VMware Cloud inventory and the networking, compute, and service capacity available in the SDDC. After the cluster has been created, Kubernetes-compatible tools can be used to interact with workloads and cluster resources, but the VMware Cloud Console is the relevant deployment interface described by the question. VMware's Tanzu deployment reference architecture for VMware Cloud on AWS discusses the deployment and infrastructure considerations for Tanzu environments at [VMware Tanzu for Kubernetes Operations documentation](#). VMware Cloud service information is also available at [VMware Cloud on AWS](#).