

Designing Cisco Enterprise Networks (ENSLD)

Cisco 300-420

Version Demo

Total Demo Questions: 51

Total Premium Questions: 513

Buy Premium PDF

<https://passqueen.com>

support@passqueen.com

Topic Break Down

Topic	No. of Questions
Topic 1, Advanced Addressing and Routing Solutions	195
Topic 2, Advanced Enterprise Campus Networks	100
Topic 3, WAN for Enterprise Networks	137
Topic 4, Network Services	25
Topic 5, Automation	56
Total	513

QUESTION NO: 1

Refer to the exhibit. Company A acquired Company B, and the companies want to use the same shared core. However, the companies must not learn each other's prefixes. Which two address family solutions meet these requirements? Choose two.

- A. IPv4 unicast for Company A and Company B
- B. VPNv4 for the shared core and for Company A and Company B
- C. IPv4 unicast and VPNv4 for Company A and Company B
- D. IPv4 unicast and VPNv4 for shared core

ANSWER: A D

Explanation:

IPv4 unicast for Company A and Company B is appropriate at the customer-facing edge. Each company can use standard IPv4 routing toward its connected provider-edge router, while the provider edge associates received routes with that company's dedicated VRF. The VRF provides the separate customer routing context required to prevent route visibility between the acquired companies.

IPv4 unicast and VPNv4 for shared core is also required in an MPLS Layer 3 VPN design. The provider core uses its IPv4 unicast routing infrastructure to provide reachability between provider-edge loopbacks and to support MPLS transport. Provider-edge routers exchange customer VPN routes using Multiprotocol BGP VPNv4. VPNv4 adds a route distinguisher to an IPv4 prefix, making prefixes from distinct VRFs unique even when address space overlaps. Route targets govern import and export policy, so Company A routes remain in Company A's VRF and Company B routes remain in Company B's VRF. This combination enables a common backbone while maintaining independent routing domains and controlled prefix distribution. Cisco's MPLS VPN documentation describes VRFs, route distinguishers, and route targets as the foundation for this customer-route separation. See [Cisco MPLS Layer 3 VPN Configuration Guide](#) and [Cisco BGP MPLS Layer 3 VPN documentation](#).

QUESTION NO: 2

Which design achieves SD-WAN control plane redundancy?

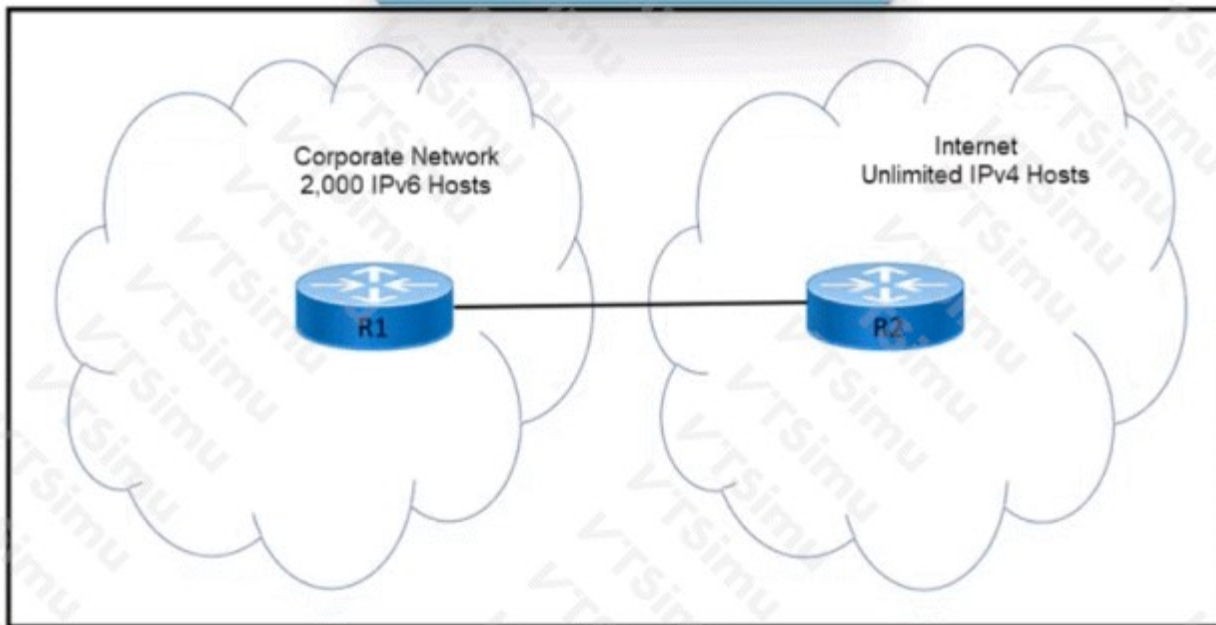
- A. Configuring BFD on the WAN Edge routers
- B. Using multiple vSmart controller instances
- C. Deploying using a virtual platform like UCS or CSP
- D. Managing the underlay network with OMP

ANSWER: B

Explanation:

Using multiple vSmart controller instances achieves Cisco SD-WAN control-plane redundancy because vSmart controllers form the centralized control plane that distributes Overlay Management Protocol (OMP) route, transport-location (TLOC), service, and policy information to WAN Edge devices. WAN Edges can maintain control connections to more than one vSmart controller. This lets the overlay continue to receive control-plane information and policy if an individual vSmart controller becomes unavailable. For a resilient design, vSmart controllers should be deployed as separate instances and placed across independent failure domains, such as different availability zones or data centers, with adequate WAN transport diversity. The controller certificate and control-connection mechanisms allow WAN Edges to establish and use these redundant controller relationships automatically. Cisco documents that the vSmart controller is responsible for control-plane functions and that WAN Edge routers establish secure control connections with controllers. Cisco's control-plane overview is available in [Cisco SD-WAN Control Plane Overview](#). Additional controller-role and deployment guidance is provided in the [Cisco SD-WAN Installation and Upgrade Guide](#).

QUESTION NO: 3



Refer to the exhibit. An engineer must design an address translation solution to provide Internet connectivity for the corporate network. The design is restricted to the 172.16.168.0/22 subnet. Which solution must the engineer choose?

- A. stateful NAT64
- B. stateless NAT64
- C. stateful NAT66
- D. stateless NAT66

ANSWER: A

Explanation:

Stateful NAT64 is the appropriate solution because it provides translation between IPv6-only corporate endpoints and IPv4 destinations on the Internet while conserving the available IPv4 address space. The 172.16.168.0/22 network contains a limited number of IPv4 addresses, so the translation device should maintain per-flow state and dynamically map many IPv6 client sessions to a smaller pool of IPv4 source addresses. Port translation allows multiple simultaneous client connections to share an IPv4 address, making the restricted IPv4 subnet usable for a larger IPv6 corporate population. Stateful NAT64 also performs the required IPv6-to-IPv4 protocol and address-family translation for sessions initiated by IPv6 clients toward IPv4-only services. Cisco documents NAT64 stateful translation as maintaining a binding table that maps IPv6 transport addresses to IPv4 transport addresses, enabling address and port multiplexing. This behavior is designed for IPv6-to-IPv4 connectivity where the IPv4 side has constrained addressing resources. See Cisco's [NAT64 configuration overview](#) and the IETF's [RFC 6146 Stateful NAT64 specification](#).

QUESTION NO: 4

A customer plans to adopt distributed QoS in their enterprise WAN. The policy must allow for individual packet marking according to the type of treatment required and for forwarding based on hop-by-hop treatment locally defined on each device. Which technology must the customer select?

- A. CBWFQ
- B. LLQ
- C. Diffserv

ANSWER: C

Explanation:

Diffserv is correct because it is a scalable, distributed QoS architecture in which traffic is classified and marked at network boundaries, typically by setting the IP Differentiated Services Code Point (DSCP). The DSCP marking identifies the per-hop behavior (PHB) that the packet should receive, such as expedited forwarding or assured forwarding. Each network device then applies its locally configured QoS policy to packets based on that marking, providing the appropriate queuing, scheduling, policing, or drop treatment at every hop. This model does not require every router to maintain per-flow reservation state, which makes it well suited to an enterprise WAN that needs consistent treatment across many devices and links. Cisco describes DiffServ as using DSCP markings to select a PHB at each network node, while RFC 2474 defines the differentiated-services field and its role in selecting forwarding behavior. See [Cisco IOS DSCP QoS Overview](#) and [RFC 2474](#).

QUESTION NO: 5

Which two overlay network design considerations must be made for a Cisco SD-Access network? (Choose two.)

- A. LAN automation for deployment
- B. Layer 3 to the access design
- C. Reduce subnets and simplify DHCP management
- D. Dedicated IGP process for the fabric
- E. Avoid overlapping IP subnets

ANSWER: C E

Explanation:

Reduce subnets and simplify DHCP management is a key Cisco SD-Access overlay design consideration because the fabric decouples endpoint addressing from a specific physical access location. A larger, appropriately scoped IP pool can be extended across fabric edge nodes in the same virtual network, enabling endpoint mobility without requiring a client to obtain a new address each time it moves. This reduces the number of DHCP scopes that must be created, maintained, monitored, and synchronized, while retaining clear address-pool boundaries based on policy and virtual-network requirements.

Avoid overlapping IP subnets is also essential for predictable endpoint reachability and operations within a virtual network. Unique endpoint address space allows the fabric control plane to maintain unambiguous endpoint-to-location mappings and permits DHCP, DNS, routing, troubleshooting, security policy, and external connectivity to operate consistently. Address planning should therefore reserve sufficiently sized, nonoverlapping pools for the relevant SD-Access virtual networks and deployment domains before endpoints are onboarded. Cisco's SD-Access design guidance discusses simplifying subnet and DHCP design and planning address space for the fabric overlay. See the [Cisco SD-Access Design Guide](#) and [Cisco SD-Access overview](#).

QUESTION NO: 6

What are two advantages of the Cisco SD-WAN technology 9 (Choose two)

- A. Improved application experience
- B. Easier deployment
- C. Optimized cloud connectivity
- D. Proactive network management
- E. Consistent connectivity

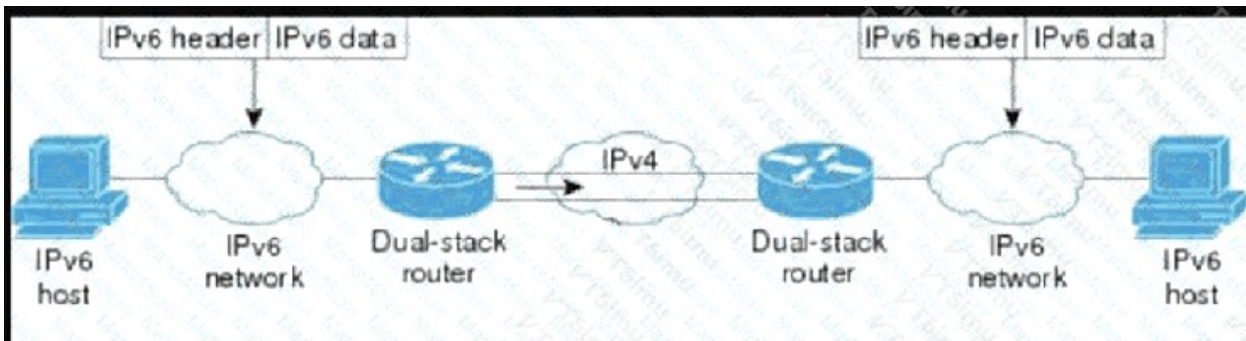
ANSWER: A C

Explanation:

Cisco SD-WAN provides an **Improved application experience** by identifying applications and continuously measuring the performance characteristics of every available WAN transport, including loss, latency, and jitter. Centralized policy can then steer business-critical traffic over the path that meets its configured service-level requirements. This application-aware routing lets the WAN react to changing path conditions and helps maintain the quality required by real-time and important business applications. Cisco also supports techniques such as forward error correction and packet duplication for selected application traffic, further protecting the user experience when a path is impaired.

Optimized cloud connectivity is a core SD-WAN advantage because branches can use policy-controlled direct access to internet, SaaS, and IaaS resources rather than sending all cloud-bound traffic through a data center. Cisco SD-WAN Cloud OnRamp capabilities extend this model by simplifying connectivity and improving path selection for cloud applications and workloads. Together, these functions make cloud access more efficient while retaining centrally managed policy and visibility. Cisco describes application optimization and cloud-optimized connectivity as key SD-WAN capabilities in its [SD-WAN overview](#) and [Cloud OnRamp documentation](#).

QUESTION NO: 7



Refer to the exhibit. A company architect proposed this network design as a part of the IPv6 migration strategy. What are two advantages of this design? Choose two.

- A. It permits any number of devices to join the overlay network seamlessly.
- B. It establishes an independent network that does not share fate with the underlay.
- C. It allows multiple independent networks to be built on top of a shared underlay.
- D. It provides increased scalability without increasing forwarding overhead.
- E. It enables the transport of protocols that are unsupported by the underlay.

ANSWER: C E

Explanation:

This design uses an overlay tunnel to carry IPv6 traffic over an existing underlay that may provide only IPv4 transport. An overlay separates logical connectivity from the physical transport infrastructure: tunnel endpoints encapsulate the passenger packet for transit through the underlay, then decapsulate it at the far endpoint. This enables multiple independent logical networks to coexist over the same shared physical and routed underlay. Each overlay can have its own addressing, routing policies, segmentation, and traffic handling while using the common transport network.

The design also enables protocols that the underlay does not natively support to traverse it. In an IPv6 migration, IPv6 packets can be encapsulated for transport across IPv4-only infrastructure, allowing IPv6-capable sites to communicate without requiring every intermediate underlay device to be upgraded to IPv6 at the same time. This supports phased migration while preserving the existing transport environment. Cisco documents that tunneling encapsulates IPv6 packets within IPv4 packets so they can traverse IPv4 infrastructure, and its enterprise design guidance describes overlays as logical networks built over an underlay. See [Cisco: IPv6 Tunneling](#) and [Cisco: Software-Defined Access](#).

QUESTION NO: 8

An architect is redesigning a campus network from a Layer 2 access design to a routed access design. Which two benefits does a Layer 3 access design provide? (Choose two.)

- A. It eliminates spanning-tree-blocked uplinks.
- B. It reduces the scope of Layer 2 failures.
- C. It extends VLANs across all access switches.
- D. It removes the need for first-hop redundancy protocols.
- E. It allows hosts to use the distribution switch as a Layer 2 default gateway.

ANSWER: A B

Explanation:

It eliminates spanning-tree-blocked uplinks because access-to-distribution connections are routed links rather than Layer 2 trunks. Equal-cost paths can be used actively by the routing protocol instead of leaving redundant Layer 2 links in a blocking state.

It reduces the scope of Layer 2 failures because each access switch operates as a separate Layer 3 boundary. A spanning-tree event, broadcast issue, or VLAN-related failure is contained to the local access segment rather than affecting a large Layer 2 domain. Cisco campus design guidance describes routed access as a design that improves resilience and limits Layer 2 failure domains. See the [Cisco Campus Network Design overview](#).

QUESTION NO: 9

Prior to establishing full-mesh IPsec tunnels in a typical Cisco SD-WAN deployment, which mechanism do WAN Edge routers use to exchange Key information for data plane encryption?

- A. They use vSmart controllers as key exchange servers.
- B. They use vManage as a key exchange server.
- C. They use IKEv2 when exchanging keys with each other.
- D. They use vBond as a key exchange server.

ANSWER: A

Explanation:

They use vSmart controllers as key exchange servers. In Cisco SD-WAN, WAN Edge routers first establish secure control-plane connections to the controller infrastructure. The vSmart controller acts as the centralized key-management component for the SD-WAN overlay: it generates and distributes the keying material needed by WAN Edge routers to build secure IPsec data-plane tunnels to one another. This process supports the scalable any-to-any, full-mesh nature of the SD-WAN fabric without requiring each pair of WAN Edge routers to perform a separate traditional peer-to-peer key negotiation.

After receiving the relevant key information through the secure control plane, WAN Edge routers can authenticate and encrypt their direct IPsec tunnels. The vSmart controller also coordinates key rotation so that encryption material can be refreshed across the fabric while maintaining secure connectivity. Cisco describes vSmart controllers as responsible for distributing encryption keys to WAN Edge devices and for controlling the SD-WAN overlay through OMP-based control-plane policy and route distribution. See Cisco's [SD-WAN IPsec tunnel documentation](#) and the [Cisco SD-WAN components overview](#).

QUESTION NO: 10

A network engineer must design a multicast solution to prevent the spoofing of multicast streams and ensure efficient bandwidth utilization. The network will be merged with another multicast domain in the future, and the merge must require minimum effort. Which two solutions meet the customer requirements? (Choose two.)

- A. PIM-SSM
- B. IGMPv3
- C. IGMPv2
- D. PIM-SM
- E. MSDP

ANSWER: A B

Explanation:

PIM-SSM and **IGMPv3** together provide source-specific multicast, in which a receiver requests a channel using both a multicast group and an explicitly identified source, represented as (S,G). This source-inclusive subscription model limits delivery to the approved source rather than accepting traffic for a group from arbitrary sources. It therefore provides meaningful protection against unwanted or spoofed multicast sources, while PIM builds efficient source-based shortest-path distribution trees only where receivers have requested the channel.

IGMPv3 is required at the receiver-facing edge for SSM because it adds source-filtering capabilities, including the ability for hosts to report that they want traffic for a group only from specified source addresses. PIM-SSM uses those source-specific membership reports to create the corresponding multicast forwarding state. This avoids reliance on rendezvous-point and shared-tree state, simplifying multicast operations and reducing dependencies when two routed networks are later combined. Assuming unicast reachability and consistent SSM group-range policy, the merged environment can continue to route source-specific joins without establishing an interdomain shared-tree infrastructure.

The IETF SSM architecture is defined in [RFC 4607](#), and the source-filtering behavior required from hosts and routers is specified in [RFC 3376 \(IGMPv3\)](#).

QUESTION NO: 11

Which resource is required for the vBond orchestrator to onboard a WAN Edge router via manual configuration?

- A. vSmart hostname
- B. domain name
- C. NAT
- D. organization name

ANSWER: D

Explanation:

The required resource is **organization name**. In a manually configured Cisco Catalyst SD-WAN WAN Edge, the system configuration must include the organization name used by the SD-WAN overlay. This value identifies the administrative SD-WAN domain and must be consistent with the organization name configured on the controllers, including vBond. It is used as part of the device and controller authentication process when the WAN Edge establishes its initial secure control connection to vBond.

During onboarding, the WAN Edge reaches vBond using its configured vBond information and presents its identity and certificate-based credentials. vBond validates that the joining device belongs to the same SD-WAN organization before orchestrating controller discovery and assisting the device in forming control-plane connections. A matching organization name is therefore a fundamental system-identity parameter for successful manual onboarding and control-plane establishment. Cisco documentation lists `organization-name` as a system-level configuration element for Cisco SD-WAN devices and describes the role of vBond in authenticating and orchestrating WAN Edge connections. See Cisco's [Configure System Settings](#) documentation and the [vBond Orchestrator configuration documentation](#).

QUESTION NO: 12

Refer to the exhibit. An architect with an employee ID: 4542:60:170 is designing a campus Layer 2 infrastructure. The design requires a PoE power budget that varies from 30-60 W. In addition, power must be provided continuously to some endpoints and must be supported even during the reloading of edge switches. Which solution must the architect select?

- A. PoE Plus
- B. Fast PoE
- C. Universal PoE
- D. Perpetual PoE
- E. Cisco Universal PoE with Perpetual PoE

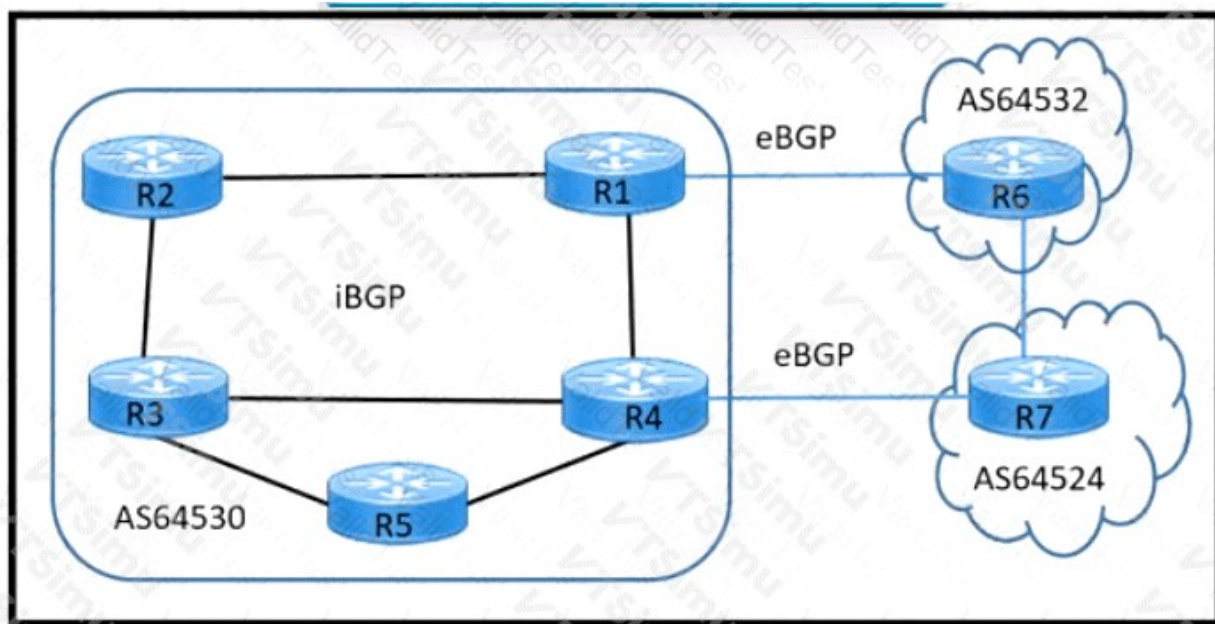
ANSWER: E

Explanation:

Cisco Universal PoE (UPOE) with Perpetual PoE is required because the design has two separate power requirements. Cisco UPOE supplies up to 60 W of power per port, meeting the upper end of the stated 30–60 W endpoint requirement. It is intended for higher-power devices such as advanced wireless access points, thin clients, and other endpoints that need more power than conventional PoE+ can deliver. UPOE can also supply lower power levels as negotiated with the attached powered device, so it accommodates endpoints across the required range.

Perpetual PoE addresses the continuity requirement. It preserves PoE output during a switch reload, so powered endpoints that are configured for this capability do not lose power while the edge switch control plane restarts. This is particularly important for services such as IP telephony, wireless access, cameras, and building-control endpoints where an avoidable power interruption is undesirable. Therefore, the design must combine UPOE for the required per-port power capacity with Perpetual PoE for continuous power through reload events. Cisco documents UPOE capabilities in its [Cisco UPOE technology overview](#) and describes PoE behavior and configuration, including perpetual PoE support, in the [Catalyst 9300 PoE configuration guide](#).

QUESTION NO: 13



Refer to the exhibit. A network engineer must design a BGP solution based on:

The route reflector must have one or more direct physical connections to the core routers (R3 and R4).

The route reflector must have full redundancy and avoid a single point of failure.

R2 to R1 link utilization is 90%. and the remaining links are less than 50% utilized.

Which two solutions must the design Include? (Choose two.)

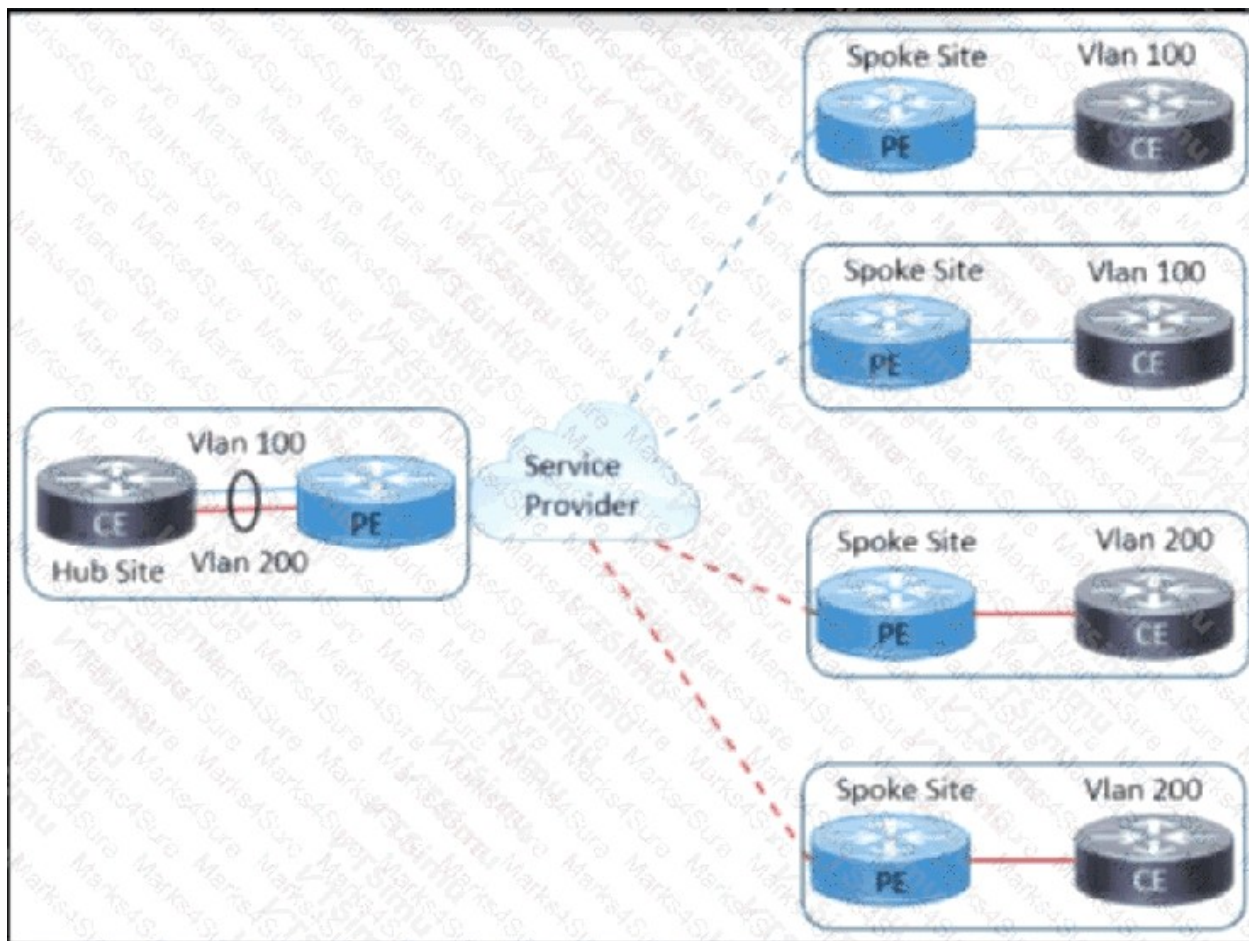
- A. Configure R1 to be a client of R2 and R4.
- B. Configure R2 to be a client of R1 and R4.
- C. Configure R3 to be a client of R2 and R4.
- D. Configure R4 to be a client of R1 and R3.
- E. Configure R5 to be a client of R3 and R4.

ANSWER: B E

Explanation:

Configure R2 to be a client of R1 and R4, and configure R5 to be a client of R3 and R4. Together, these relationships provide redundant route-reflection service: R2 has two route reflectors, R1 and R4, while R5 has two route reflectors, R3 and R4. Therefore, loss of a single reflector or a single relevant adjacency does not remove route-reflector reachability for either client. The selected reflectors also satisfy the placement requirement in the exhibit because the reflector design includes direct physical connectivity to the core routers R3 and R4. Using R3 and R4 in the route-reflector arrangement places reflection capability at the core, while R1 provides an additional reflector path for R2 without making the heavily utilized R2-to-R1 link the sole dependency. BGP route reflection reduces the need for a full mesh of internal BGP peerings; route reflectors advertise eligible paths between their clients and nonclients according to the route-reflection rules. Redundant reflectors are a standard design practice to preserve control-plane availability when a reflector fails. Cisco documents route-reflector operation and client peering behavior in its [BGP Route Reflector Configuration and Operation guide](#) and in the [Cisco IOS XE BGP Route Reflector documentation](#).

QUESTION NO: 14



Refer to the exhibit. An architect working for a service provider with an employee ID: 4763:44:876 must design a Layer 2 VPN solution that supports:

transparency of service provider devices

direct communication between CE routers attached to the same VLAN

Which solution must the design include?

- A. multiple VPWS
- B. single VPLS
- C. single VPWS
- D. multiple VPLS

ANSWER: B

Explanation:

single VPLS is the appropriate design because Virtual Private LAN Service provides a multipoint Layer 2 Ethernet service. A single VPLS instance represents one virtual bridged LAN, or broadcast domain, spanning the provider network. Each attached customer edge interface participates in that same emulated LAN and can therefore exchange Ethernet frames directly with other customer edge routers in the same VLAN. This meets the requirement for direct, any-to-any communication without requiring the customer to see or interact with the provider's internal topology.

VPLS uses provider-edge devices to learn customer MAC addresses and forward frames across pseudowires, while presenting the service to customers as a transparent Ethernet switching domain. The provider MPLS/IP transport and intermediate provider devices remain outside the customer's Layer 2 view. This transparent LAN behavior is specifically the service model defined for VPLS: it delivers a single logical LAN across geographically separated attachment circuits. Cisco's Layer 2 VPN documentation describes VPLS as a multipoint Ethernet service, and RFC 4762 defines the architecture and forwarding behavior for a virtual private LAN service. See [Cisco IOS XE VPLS documentation](#) and [RFC 4762](#).

QUESTION NO: 15

Which design consideration must be made when using IPv6 overlay tunnels?

- A. Overlay tunnels that connect isolated IPv6 networks can be considered a final IPv6 network architecture.
- B. Overlay tunnels should only be considered as a transition technique toward a permanent solution.
- C. Overlay tunnels can be configured only between border devices and require only the IPv6 protocol stack.
- D. Overlay tunneling encapsulates IPv4 packets in IPv6 packets for delivery across an IPv6 infrastructure.

ANSWER: B

Explanation:

Overlay tunnels should only be considered as a transition technique toward a permanent solution. An IPv6 overlay tunnel is commonly used to transport IPv6 packets across an IPv4-only portion of a network, allowing separately deployed IPv6-capable sites, sometimes called IPv6 islands, to communicate before the intervening infrastructure supports native IPv6 forwarding. This can be a practical migration measure when upgrading every transit device or service at once is not feasible.

The design must account for the fact that tunneling adds an encapsulation header and therefore reduces the effective packet size available to applications. The tunnel also depends on the reachability, MTU, QoS behavior, and resiliency of its IPv4 underlay. These dependencies can complicate operations, fault isolation, and path-MTU handling as the network grows. For those reasons, a long-term enterprise IPv6 design should plan to deploy native IPv6 forwarding where possible, often alongside IPv4 in a dual-stack migration model, rather than make transition tunnels the enduring architecture. Cisco identifies tunneling as one of the IPv6 transition mechanisms used where native IPv6 connectivity is not yet available; see [Cisco IPv6 solutions](#) and the [Cisco IOS XE IPv6 tunneling configuration guide](#).

QUESTION NO: 16

An engineer must design a VPN solution for a company that has multiple branches connecting to a main office. What are two advantages of using DMVPN instead of IPsec tunnels to accomplish this task?

(Choose two.)

- A. support for AES 256-bit encryption
- B. greater scalability
- C. support for anycast gateway
- D. lower traffic overhead
- E. dynamic spoke-to-spoke tunnels

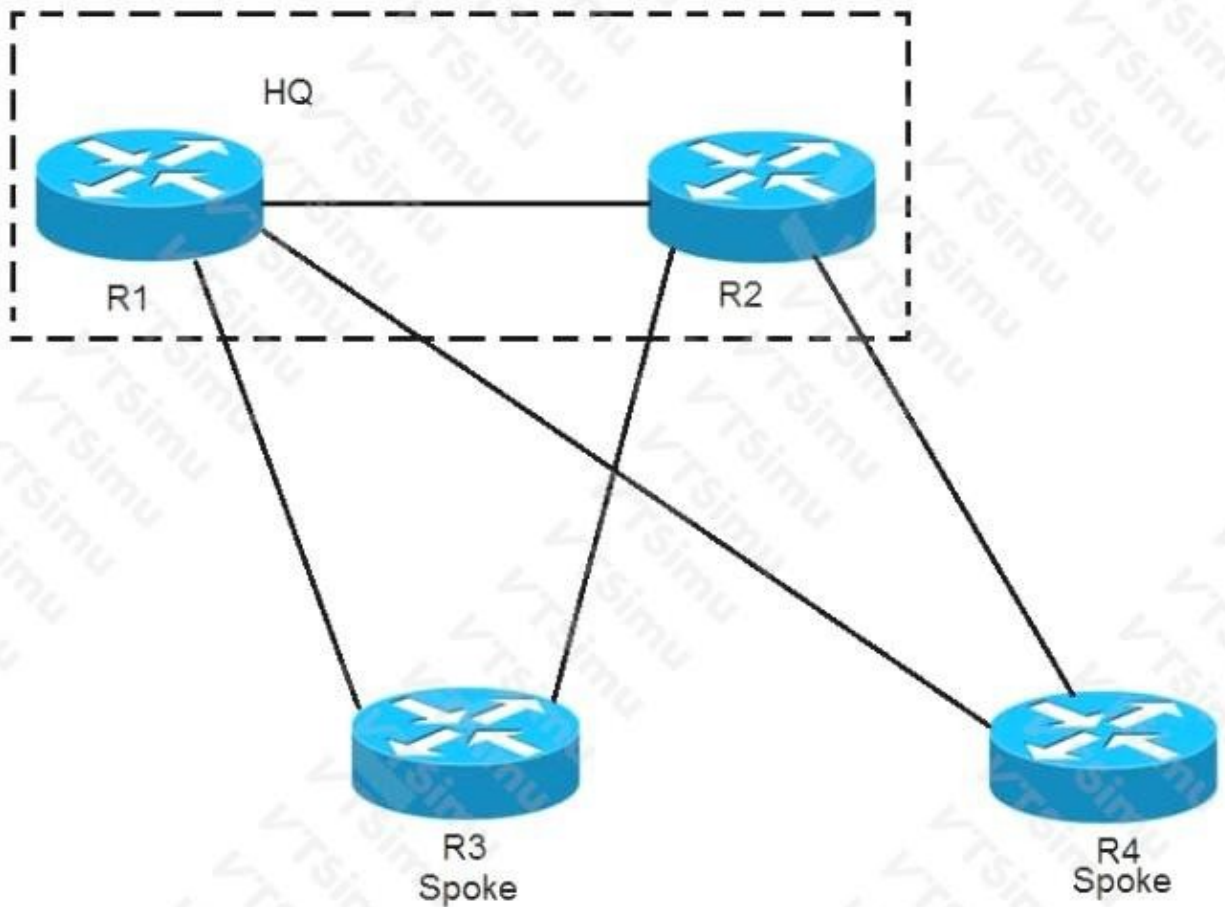
ANSWER: B E

Explanation:

greater scalability and **dynamic spoke-to-spoke tunnels** are the key advantages of Dynamic Multipoint VPN (DMVPN) for a hub-and-spoke branch design. DMVPN uses multipoint GRE interfaces, Next Hop Resolution Protocol (NHRP), and IPsec protection to let spokes register their dynamically learned addresses with a hub. This avoids configuring and maintaining a separate static GRE/IPsec tunnel interface and peer relationship for every branch-to-hub or branch-to-branch path. As the organization adds branches, the operational effort grows far more efficiently than with a full set of individually defined point-to-point IPsec tunnels.

DMVPN also enables spokes to establish direct, on-demand protected tunnels to other spokes after resolving the remote tunnel endpoint through NHRP. Consequently, traffic between branches can follow a direct path rather than being forced through the main-office hub. Direct spoke-to-spoke connectivity reduces hub transit load and can improve the path taken by interbranch traffic, while retaining centralized hub-based registration and control. Cisco describes DMVPN as a solution that combines GRE, NHRP, and IPsec to provide scalable, dynamic VPN connectivity. See Cisco's [DMVPN overview](#) and [Branch WAN design guidance](#).

QUESTION NO: 17



Refer to the exhibit. EIGRP has been configured on all links. The spoke nodes have been configured as EIGRP stubs, and the WAN links to R3 have higher bandwidth and lower delay than the WAN links to R4. When a link failure occurs at the R1-R2 link, what happens to traffic on R1 that is destined for a subnet attached to R2?

- A. R1 has no route to R2 and drops the traffic
- B. R1 load-balances across the paths through R3 and R4 to reach R2
- C. R1 forwards the traffic to R3, but R3 drops the traffic
- D. R1 forwards the traffic to R3 in order to reach R2

ANSWER: D

Explanation:

R1 forwards the traffic to R3 in order to reach R2. With the direct R1-R2 connection unavailable, EIGRP converges on the remaining reachable path to R2's attached subnet. The topology provides alternate WAN connectivity through both R3 and R4, but the route through R3 has the better EIGRP composite metric because its WAN links have higher bandwidth and lower delay. By default, EIGRP calculates its metric using bandwidth and delay, so the path through R3 is preferred over the corresponding path through R4.

The EIGRP stub configuration on the spoke routers does not prevent the hubs from advertising reachable spoke-connected networks to other routers. Stub operation is intended to limit queries and prevent a spoke from being selected as a transit path between hubs or other peers. R3 can therefore provide R1 with a valid route toward R2 and forward the traffic over its direct connectivity to R2. Cisco documents the purpose and route-advertisement behavior of EIGRP stub routing in its [EIGRP Stub Routing overview](#); EIGRP metric behavior is described in Cisco's [EIGRP technical documentation](#).

QUESTION NO: 18 - (DRAG DROP)

DRAG DROP

Drag and drop the descriptions from the left onto the corresponding VPN types on the right.

Select and Place:

Answer Area

The service provider participates in routing with the customer.

The customer controls the IP routing and policy governance.

Sites appear to each other to be directly connected at Layer 3.

Sites appear to be connected via the MPLS service provider network.

The customer initiates Layer 3 connectivity with the remote sites.

The customer establishes Layer 3 connectivity with the service provider edge device.

Layer 2 VPN

MPLS Layer 3 VPN

ANSWER:

Answer Area

The service provider participates in routing with the customer.

The customer controls the IP routing and policy governance.

Sites appear to each other to be directly connected at Layer 3.

Sites appear to be connected via the MPLS service provider network.

The customer initiates Layer 3 connectivity with the remote sites.

The customer establishes Layer 3 connectivity with the service provider edge device.

Layer 2 VPN

The customer controls the IP routing and policy governance.

Sites appear to each other to be directly connected at Layer 3.

The customer initiates Layer 3 connectivity with the remote sites.

MPLS Layer 3 VPN

The service provider participates in routing with the customer.

Sites appear to be connected via the MPLS service provider network.

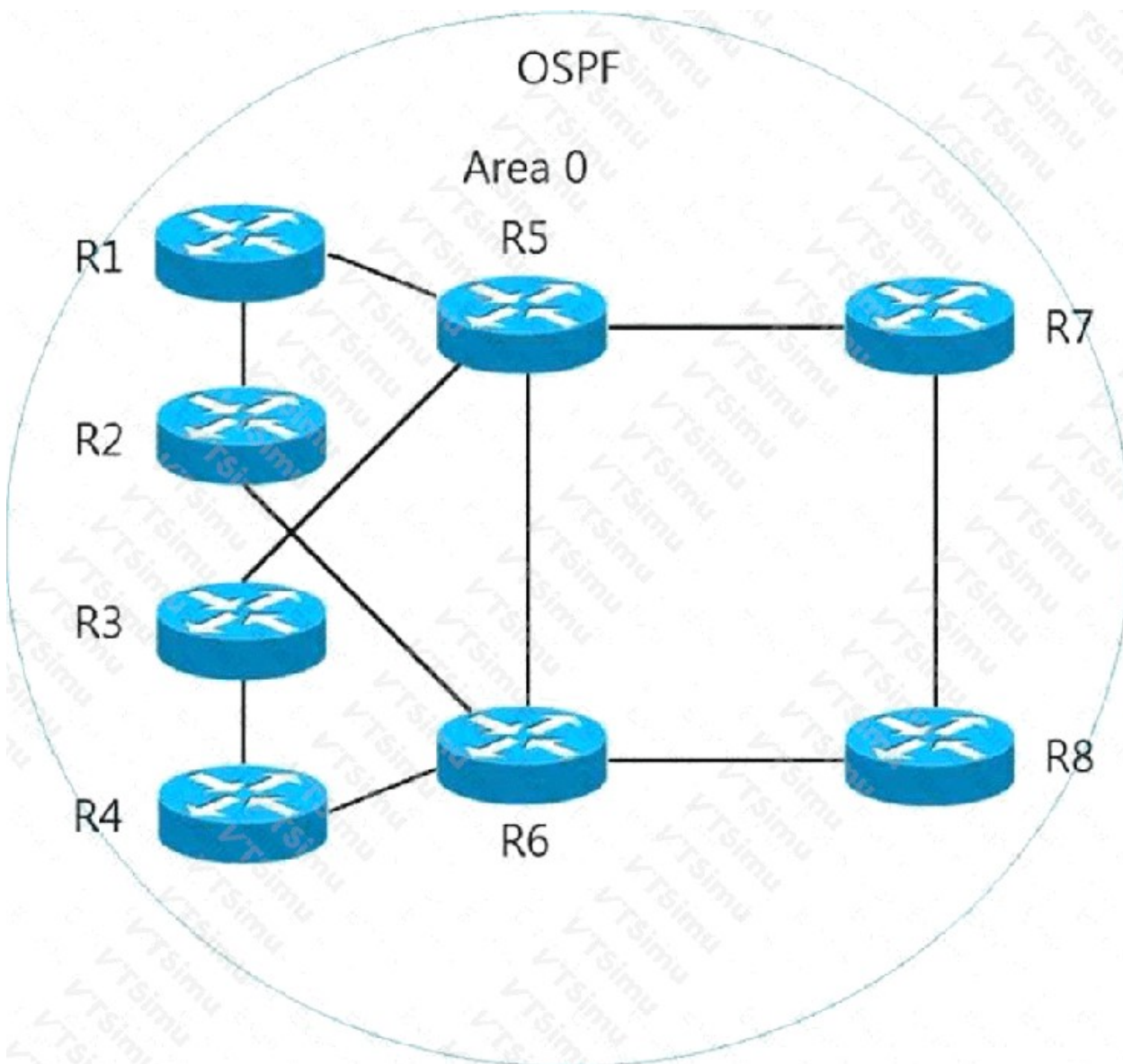
The customer establishes Layer 3 connectivity with the service provider edge device.

Explanation:

A Layer 2 VPN delivers a Layer 2 transport service between customer sites. Because the provider is extending the customer's Layer 2 domain rather than providing the customer's routed VPN service, the customer retains control of IP routing, routing policy, and the Layer 3 design. Customer routers form their own Layer 3 connectivity across the supplied Layer 2 path, toward the remote customer sites. From the customer routing perspective, the remote locations can therefore appear directly connected at Layer 3, even though the underlying transport may cross a provider network. Cisco describes Layer 2 VPN technologies as services that transport customer Layer 2 frames between locations, leaving customer routing functions with the customer. See Cisco's [Layer 2 VPN configuration documentation](#).

An MPLS Layer 3 VPN is a routed VPN service. The customer edge device establishes Layer 3 connectivity with a provider edge device, commonly through a routing relationship between the CE and PE. The provider edge device maintains the VPN routing information and participates in the routing service used to reach the customer's other VPN sites. Consequently, the customer sites appear to be connected through the MPLS service-provider network rather than through a customer-managed direct Layer 3 adjacency. MPLS VPN route separation and forwarding are implemented by provider edge devices using VPN routing and forwarding instances and MPLS transport across the provider core. Cisco's [MPLS Layer 3 VPN documentation](#) covers the PE-based VPN routing model.

QUESTION NO: 19



Refer to the exhibit. All routers currently reside in OSPF area 0. The network manager recently used R1 and R2 as aggregation routers for remote branch locations and R3 and R4 for aggregation routers for remote office locations. The network has since been suffering from outages, which are causing frequent SPF runs. To enhance stability and introduce

areas to the OSPF network with the minimal number of ABRs possible, which two solutions should the network manager recommend? (Choose two.)

- A. a new OSPF area for R1 and R2 connections, with R1 and R2 as ABRs
- B. a new OSPF area for R3 and R4 connections, with R5 and R6 as ABRs
- C. a new OSPF area for R3 and R4 connections, with R3 and R4 as ABRs
- D. a new OSPF area for R1, R2, R3, and R4 connections, with R1, R2, R3, and R4 as ABRs
- E. a new OSPF area for R1 and R2 connections, with R5 and R6 as ABRs

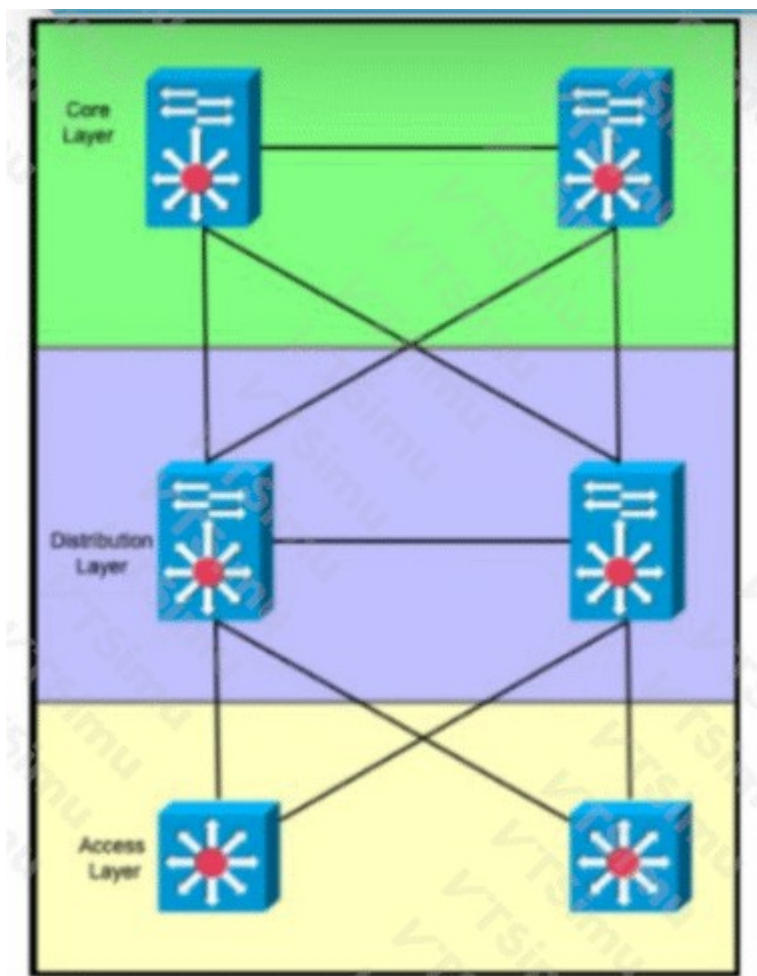
ANSWER: B E

Explanation:

The recommended design is to create a separate OSPF area for the connections aggregated by R1 and R2, with R5 and R6 providing the area-boundary role, and another separate OSPF area for the connections aggregated by R3 and R4, again with R5 and R6 as the area-boundary routers. This moves the area boundaries to the higher aggregation layer, allowing the same two routers to connect each newly created nonbackbone area to area 0. Therefore, the design introduces only two ABRs while still separating the branch and remote-office failure domains.

OSPF maintains a link-state database per area, and topology changes are flooded within the affected area. Segmenting the branch and remote-office portions into distinct areas contains link-state changes from those portions of the network, reducing the scope of SPF recalculations and improving stability in the backbone and in the other remote domain. An ABR must have an interface in the backbone and an interface in at least one nonbackbone area; R5 and R6 satisfy that function while retaining redundant paths between each new area and area 0. OSPF area-border behavior is specified in [RFC 2328](#); Cisco also describes OSPF area design and ABR operation in its [OSPF design guidance](#).

QUESTION NO: 20



Refer to the exhibit. An engineer is designing a multicampus Layer 3 Infrastructure using EIGRP as the routing protocol. The design must provide quick replies to queries in the event of a downlink, prevent unnecessary queries, and ensure that traffic does not unnecessarily transit the access layer. Which two actions must the engineer take for the network design? (Choose two.)

- A. Configure core layer switches as stub routers.
- B. Configure distribution layer switches to summarize routes to the core layer.
- C. Configure access layer switches as stub routers.
- D. Configure access layer and core layer switches as stub routers.
- E. Configure access layer switches to summarize routes to the distribution layer.

ANSWER: B C

Explanation:

Configuring distribution layer switches to summarize routes to the core layer creates route boundaries at the distribution-to-core edge. EIGRP route summarization reduces the number of individual prefixes advertised upstream and causes the summarizing router to answer queries for destinations contained within its summary. This contains EIGRP query propagation when a downstream route is lost, producing faster convergence and avoiding unnecessary query traffic through the campus backbone. The distribution layer is the appropriate aggregation point because it represents the boundary between each campus access/distribution block and the core.

Configuring access layer switches as stub routers prevents those switches from being used as transit paths between other EIGRP routers. An EIGRP stub advertises only the permitted limited route types and does not receive EIGRP queries for routes that it cannot provide, substantially reducing query scope and processing during failures. This fits the hierarchical campus role of the access layer: it should provide endpoint connectivity through its distribution-layer uplinks rather than carry traffic between other network areas. Together, distribution-layer summarization and access-layer stub configuration bound

failure domains, reduce control-plane work, and preserve intended Layer 3 traffic flow. Cisco documents EIGRP stub operation at [EIGRP Stub Routing](#) and route summarization at [EIGRP Route Summarization](#).

QUESTION NO: 21

A company has many spoke sites with two data centers. The company wants to exchange the routing information between the data centers and the spoke sites using EIGRP. All locations belong to a single AS, and auto-summarization is disabled. Which two actions must the company choose? (Choose two.)

- A. Exchange all routes between locations
- B. Summarize the routes between the hubs.
- C. Make each spoke site router a stub router
- D. Summarize the routes from spokes to the hubs.
- E. Split the network into two separate ASs

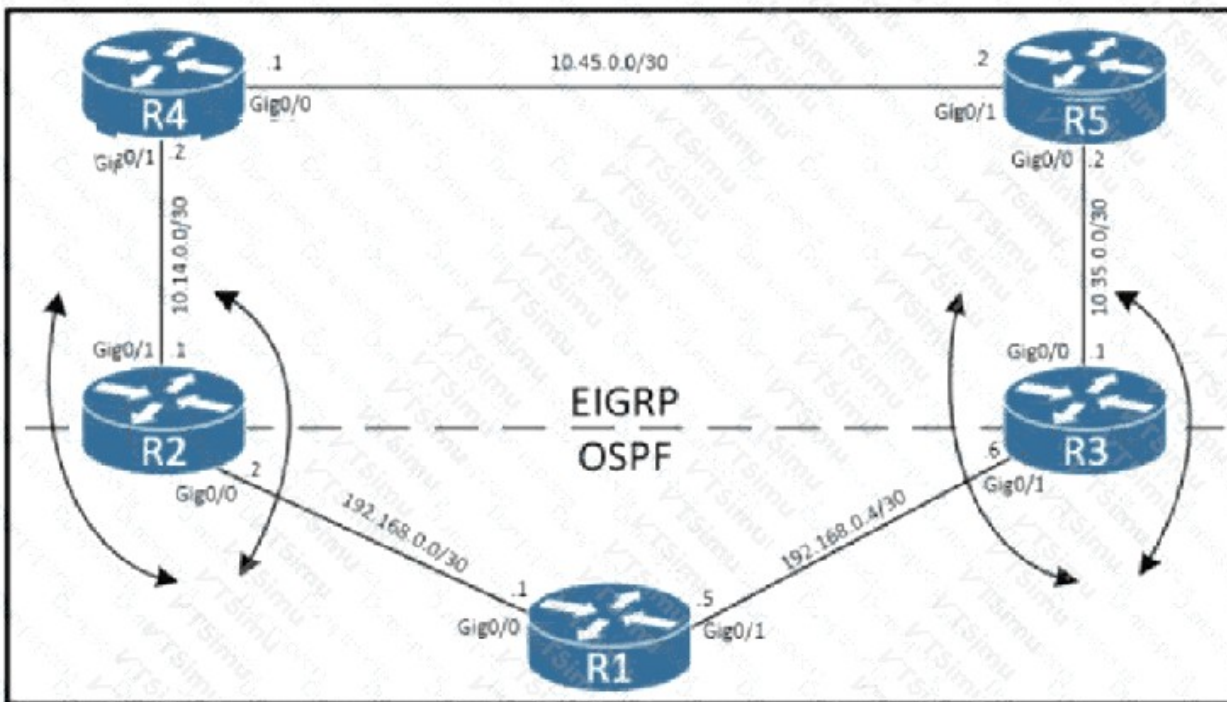
ANSWER: C D

Explanation:

Making each spoke site router a stub router is a fundamental EIGRP hub-and-spoke scalability practice. EIGRP stub routers advertise only the route categories explicitly permitted by their stub configuration, such as connected and summary routes. This tells the data-center hubs that the spoke is not a transit path and limits EIGRP query propagation into branch networks. As a result, the design reduces control-plane processing at spoke routers and improves convergence behavior as the number of sites grows. Cisco documents that EIGRP stub routing is intended for routers that should not be used as transit routers and helps constrain EIGRP query scope.

Summarizing the routes from spokes to the hubs complements stub routing. Because automatic summarization is disabled, summaries must be configured manually on the spoke-facing EIGRP interfaces. A summary aggregates the individual prefixes behind a spoke into a smaller set of advertisements sent to the data centers. This reduces routing-table entries and topology information maintained at the hubs while also limiting the scope of route changes caused by failures or updates within a spoke. The two techniques together create a scalable single-AS EIGRP design for many remote sites. See Cisco's [EIGRP configuration and design documentation](#) and the [EIGRP stub routing guide](#).

QUESTION NO: 22



Refer to the exhibit. As part of a design review of redistribution, a client requested that R2 be preferred over R3 for traffic passing toward the EIGRP domain. Which method meets this design requirement?

- A. Redistribute EIGRP into OSPF with metric-type E1 on R2 and metric-type E2 on R3.
- B. Remove the mutual redistribution on R3.
- C. Redistribute OSPF into EIGRP with metric 10000 100 255 1 1500 on R2 and metric 10 1000 255 1 1500 on R3.
- D. Redistribute EIGRP into OSPF with metric-type E2 on R2 and metric-type E1 on R3.

ANSWER: A

Explanation:

Redistribute EIGRP into OSPF with metric-type E1 on R2 and metric-type E2 on R3. This directly controls the route information advertised into the OSPF domain for destinations that reside behind the EIGRP domain. R2 and R3 act as autonomous system boundary routers when they redistribute EIGRP routes into OSPF, so OSPF routers use the external route type when selecting the preferred exit point toward those destinations.

An OSPF external type 1 route includes both the external redistribution metric and the internal OSPF cost to reach the advertising boundary router. An external type 2 route is treated as having an external metric that is independent of the internal cost to reach the boundary router, except as a tie-breaker. More importantly for this design, OSPF route selection prefers an external type 1 route over an external type 2 route for the same destination. Advertising EIGRP prefixes as E1 from R2 and as E2 from R3 therefore makes R2 the preferred OSPF exit toward the EIGRP domain while retaining R3 as an available alternate redistribution point. Cisco describes these external route types and their metric behavior in its [OSPF Design Guide](#) and [OSPF external route documentation](#).

QUESTION NO: 23

Which topology within a network underlay eliminates the need for first hop redundancy protocols while improving fault tolerance, increasing resiliency, and simplifying the network?

- A. virtualized topology
- B. routed access topology
- C. Layer 2 topology

ANSWER: B

Explanation:

routed access topology is correct because it extends Layer 3 routing to the access layer rather than relying on a Layer 2 access domain with a default gateway redundancy mechanism. Each access switch operates as a routing peer and can use equal-cost paths toward the distribution or fabric layer. Hosts retain resilient connectivity through the routed design, and traffic can be forwarded across available paths without requiring a first-hop redundancy protocol to present a shared default-gateway address.

This approach also improves failure handling because routing protocols detect and reconverge around a failed uplink or neighboring device. It simplifies the underlay by reducing dependence on large Layer 2 failure domains and associated gateway redundancy design considerations. In Cisco enterprise campus and fabric-oriented designs, a routed access underlay provides scalable, deterministic forwarding and supports fast convergence through routed links and equal-cost multipath forwarding. Cisco describes routed access as a design in which Layer 3 is extended to the access layer, reducing reliance on traditional Layer 2 mechanisms. See Cisco's [Campus Network Design overview](#) and the [Cisco Enterprise Networking solutions](#) documentation.

QUESTION NO: 24

An architect is designing a Cisco SD-WAN policy for voice traffic. The design must measure loss, latency, and jitter continuously and move voice traffic to an alternate tunnel when the preferred path no longer meets the application requirements. Which two components must the architect use? (Choose two.)

- A. BFD
- B. Application-aware routing policy
- C. Control policy
- D. TLOC extension
- E. OMP route redistribution

ANSWER: A B

Explanation:

BFD is required because Cisco SD-WAN WAN Edge devices use Bidirectional Forwarding Detection probes across active tunnels to collect path-quality measurements, including loss, latency, and jitter. These measurements provide the real-time transport health information used for application-aware path selection.

An application-aware routing policy is required to associate an application with an SLA class and define the preferred and backup transport paths. When BFD measurements show that a tunnel violates the configured SLA thresholds, the policy directs eligible traffic to another available tunnel that satisfies the SLA. Cisco describes these functions in its [Application-Aware Routing Design Guide](#).

QUESTION NO: 25



Refer to the exhibit. A customer has two eBGP peerings from a single CE router toward two service providers. The customer has hired an architect to design a solution to ensure certain traffic enters the customer's network through interface g0/0. Which solution must the architect include in the design?

- A. Advertise a lower MED value toward the less preferred service provider.
- B. Prepend additional AS on the AS path toward the preferred service provider.
- C. Break aggregated routes into longer prefixes and advertise to the preferred service provider.
- D. Set a higher local preference to the preferred service provider path.

ANSWER: C

Explanation:

Break aggregated routes into longer prefixes and advertise to the preferred service provider. This is an effective inbound traffic-engineering technique because routers apply longest-prefix matching before evaluating BGP path attributes. For example, the customer can advertise an aggregate prefix through both providers to maintain general reachability, while advertising a more-specific subnet only to the provider reached through g0/0. Remote networks that send traffic to addresses within that subnet then select the more-specific route, causing their traffic to arrive through the service provider and CE interface associated with g0/0.

This approach gives the customer a practical way to influence how external networks reach selected destinations without requiring the customer to control routing policies within either provider AS. The design must account for the fact that Internet providers may filter prefixes longer than an accepted minimum length, commonly /24 for IPv4, and the more-specific advertisement must be propagated by the preferred provider. Cisco documents that route lookup uses the longest matching prefix, and its BGP traffic-engineering guidance identifies more-specific route advertisements as a method for influencing inbound traffic. See [Cisco: BGP Path Selection Algorithm](#) and [Cisco: BGP Inbound Traffic Engineering](#).

QUESTION NO: 26

An organization is designing IPv6 addressing for a routed campus. The design must support efficient route summarization at distribution boundaries and must allow each VLAN to use a standard IPv6 LAN prefix length. Which two addressing practices should the architect use? (Choose two.)

- A. Allocate contiguous prefixes to each distribution block.
- B. Assign a /64 prefix to each VLAN.
- C. Assign a /128 prefix to each VLAN.
- D. Allocate random prefixes to access VLANs to improve security.
- E. Use a /127 prefix for every user VLAN.

ANSWER: A B

Explanation:

Allocate contiguous prefixes to each distribution block because contiguous addressing allows the distribution layer to advertise an aggregated route toward the core. Summarization reduces the number of specific prefixes in the core routing table and limits the propagation of access-layer topology changes.

Assign a /64 prefix to each VLAN because /64 is the standard IPv6 prefix length for LAN segments using Stateless Address Autoconfiguration. It provides a 64-bit interface identifier portion and supports common IPv6 host behaviors. Smaller prefixes can prevent SLAAC operation, while a consistent /64 allocation simplifies operations and address planning. IPv6 addressing architecture is defined in [RFC 4291](#), and SLAAC requirements are described in [RFC 4862](#).

QUESTION NO: 27

Refer to the exhibit. An architect must design an IP addressing scheme for a multisite network connected via a WAN transit. The campus site must accommodate 12,000 devices and the branch sites must accommodate 1,000 devices. Which address scheme optimizes network device resources, contains convergence events to the different blocks of the network, and ensures future growth of the network?

- A. •Branch1: 10.0.192.0/21
- B. •Branch1: 10.255.0.0/20
- C. •Branch1: 10.64.0.0/10
- D. •Branch1: 10.0.64.0/21

ANSWER: A

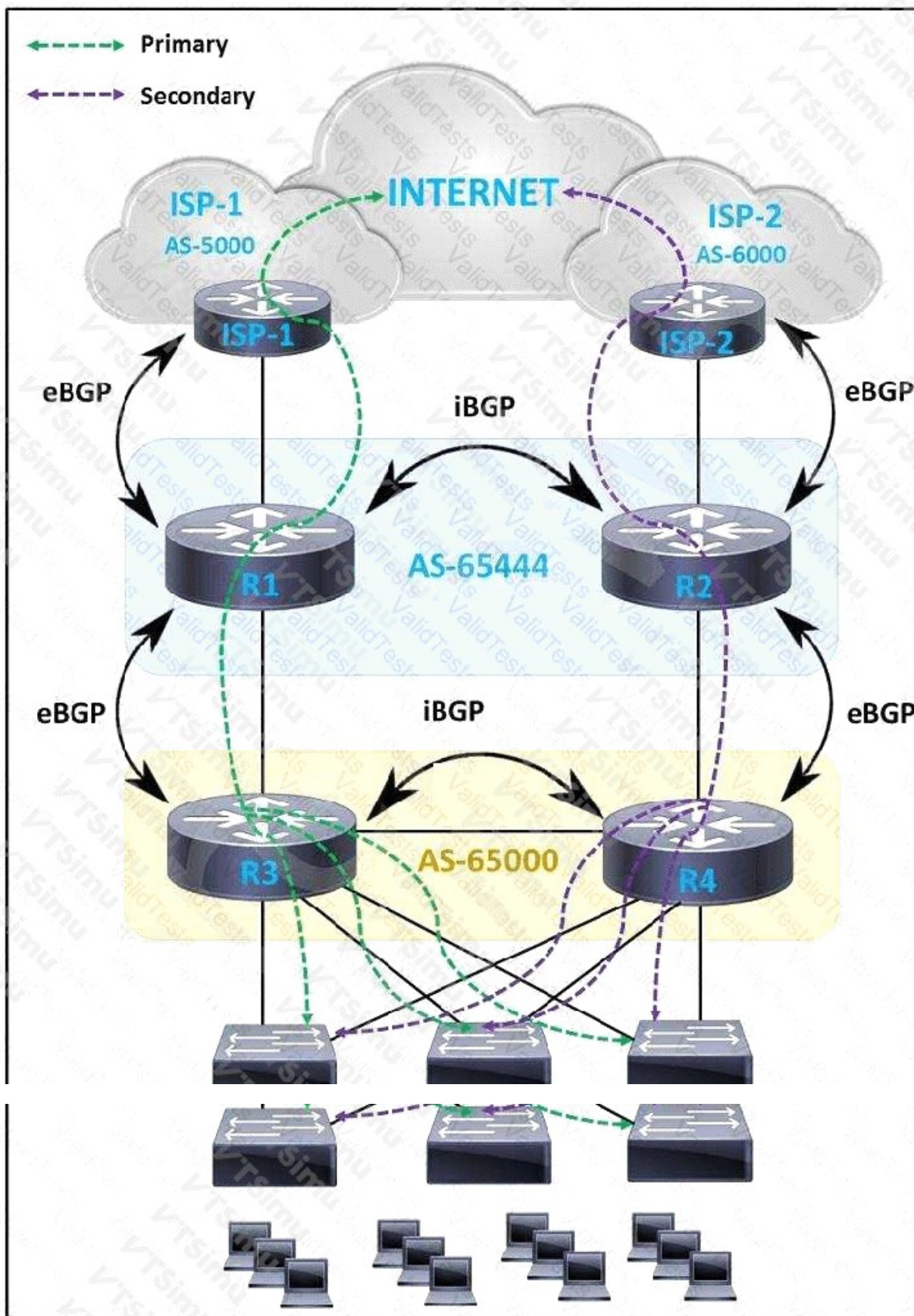
Explanation:

•Branch1: 10.0.192.0/21 is the appropriate allocation because a /21 IPv4 prefix provides 2,048 total addresses and 2,046 usable host addresses. This comfortably supports a branch requirement of 1,000 devices while avoiding unnecessary consumption of the enterprise private address block. The allocation is also aligned on a valid /21 boundary: the third octet increments in blocks of eight for this mask, so 10.0.192.0 is a proper network address.

A scalable enterprise plan assigns contiguous, hierarchical address blocks to sites. This lets the WAN and core advertise a summarized route for the branch rather than carrying routes for every internal subnet. Route summarization reduces routing-table size and limits the scope of topology changes and convergence events, because instability within a site can remain contained behind its summary. Placing Branch1 in a distinct, appropriately sized block also leaves adjacent address space available for additional branch subnets or growth without requiring renumbering. The 10.0.0.0/8 range is private IPv4 space suitable for internal enterprise addressing, as defined by [RFC 1918](#). Cisco's discussion of hierarchical network design likewise emphasizes modularity and limiting the impact of failures and changes through clear network boundaries: [Cisco Borderless Network Design Guide](#).

QUESTION NO: 28

Refer to the exhibit.



An engineer must design a WAN solution so that ISP-1 is always preferred over ISP-2. The path via ISP-2 is considered as a backup and must be used only when the path to ISP-1 is down. Which solution must the engineer choose?

A. - Routes advertised to ISP-1: 0x AS-path prepend

- B. - Routes advertised to ISP-1: 0x AS-path prepend; routes advertised to ISP-2: 5x AS-path prepend
- C. - Routes advertised to ISP-1: 0x AS-path prepend
- D. - Routes advertised to ISP-1: 5x AS-path prepend

ANSWER: B

Explanation:

Routes advertised to ISP-1: 0x AS-path prepend; routes advertised to ISP-2: 5x AS-path prepend is the appropriate design for influencing inbound traffic toward the enterprise. BGP normally prefers a route with a shorter AS-path after higher-priority attributes such as weight and local preference have been evaluated. Advertising the enterprise prefixes to ISP-1 without adding the enterprise AS preserves the shortest available path through ISP-1. Advertising the same prefixes to ISP-2 with five copies of the enterprise AS makes that route substantially less attractive to external BGP speakers that receive both advertisements.

As a result, remote networks select the path through ISP-1 under normal operating conditions. If the ISP-1 path or its advertisements are withdrawn because of a failure, the longer path learned through ISP-2 remains available and is selected as the backup path. This approach is a standard BGP traffic-engineering technique when the organization needs to influence how external networks enter its WAN. Cisco documents AS-path length as a BGP best-path selection criterion and describes AS-path prepending as a method for influencing inbound routing: [Cisco: BGP Path Selection Algorithm](#) and [Cisco: BGP Best Path Selection Algorithm](#).

QUESTION NO: 29

How does a model-driven telemetry dial-out approach function?

- A. The device initiates a session to the collector based on the subscription.
- B. The collector initiates a session to the device and subscribes to data to be streamed.
- C. The collector Initiates a session to the device and gets the data of a previously defined subscription.
- D. The device initiates a session to the collector and negotiates a subscription.

ANSWER: A

Explanation:

The device initiates a session to the collector based on the subscription. In model-driven telemetry dial-out, the network device is configured with a static telemetry subscription that identifies the data to export, the encoding, the update policy, and the destination receiver. When that subscription is active, the device itself establishes the outbound transport session to the configured collector and pushes the selected YANG-modeled operational data at the defined periodic or on-change interval. This design is especially useful when the collector cannot initiate inbound connections to managed devices because of addressing, NAT, or security-policy constraints.

The subscription is therefore the device-side instruction that drives the export. It includes receiver details such as the destination address and port, along with the telemetry data selection and streaming characteristics. After the device establishes connectivity to the collector, it continuously exports telemetry according to that configured subscription, allowing the collector to ingest and analyze near-real-time state data without repeatedly polling individual operational commands. Cisco's model-driven telemetry documentation describes dial-out subscriptions as subscriptions configured on the device that send data to a receiver. See the [Cisco IOS XE Model-Driven Telemetry documentation](#) and the [Cisco DevNet Model-Driven Telemetry overview](#).

QUESTION NO: 30

An architect must address sustained congestion on the access and distribution uplink of network. QoS has already been implemented and optimized, but it is no longer effective in ensuring optimal network performance. Which two solutions should the architect use to improve network performance? (Choose two)

- A. Reconfigure QoS based on the IntServ model
- B. Utilize random early detection to manage queues
- C. Implement higher-speed uplink interfaces
- D. Bundle additional uplinks into logical EtherChannels
- E. Configure selective packet discard to drop noncritical network traffic.

ANSWER: C D

Explanation:

When congestion is sustained after QoS has already been optimized, the appropriate design response is to increase available uplink capacity rather than further tune traffic treatment. **Implement higher-speed uplink interfaces** directly increases the bandwidth available between the access and distribution layers. This reduces interface serialization and queue buildup, allowing the network to accommodate the ongoing traffic load while preserving the QoS policy's differentiated treatment for business-critical applications.

Bundle additional uplinks into logical EtherChannels also increases aggregate bandwidth and provides link resiliency. EtherChannel presents multiple physical links as one logical port channel to Layer 2 and Layer 3 protocols, allowing traffic flows to be load-balanced across member links without creating a spanning-tree-blocked redundant path. This is particularly useful when replacing an uplink with a faster interface is not immediately feasible, or when additional physical links are available. Capacity planning should account for the hashing-based distribution of flows across EtherChannel members, but the port channel remains an established campus-design method for scaling access-to-distribution connectivity. Cisco's QoS overview is available at [Cisco Quality of Service](#), and Cisco documents EtherChannel operation at [Cisco EtherChannel Configuration Guide](#).

QUESTION NO: 31

What is the function of the multicast Reverse Path Forwarding check?

- A. It allows for a loop-free distribution tree from the source to receivers.
- B. It serves as an Auto RP Mapping agent.
- C. It prevents bootstrap messages from reaching all routers.
- D. It is used to discover and announce RP-set information.

ANSWER: A

Explanation:

It allows for a loop-free distribution tree from the source to receivers. The multicast Reverse Path Forwarding (RPF) check is the fundamental loop-prevention mechanism used when forwarding multicast traffic. When a router receives a multicast packet, it checks whether that packet arrived on the interface that the router would use to reach the packet's source, according to its unicast routing information (or the applicable multicast routing table). If the incoming interface is the expected RPF interface, the router accepts the packet and can forward copies out appropriate downstream interfaces. This process causes multicast forwarding to follow a reverse shortest-path tree rooted at the source, producing an efficient, loop-free path toward receivers.

Because a router forwards a source's multicast traffic only when it arrives through the valid reverse path, duplicate packets that could circulate through redundant network links are discarded. The RPF check therefore enables multicast routing protocols and forwarding mechanisms to build and maintain source-to-receiver distribution trees without requiring every receiver to signal every potential path. The RPF neighbor and interface selection are derived from routing reachability to the source, making correct unicast routing an essential dependency for predictable multicast forwarding. Cisco's multicast overview describes RPF as the mechanism that prevents multicast packet loops and supports distribution-tree forwarding: [Cisco IP Multicast Technology Overview](#). See also Cisco's [IP Multicast Configuration Guide](#).

QUESTION NO: 32

A customer's environment includes hosts that support IPv6-only. Several of these hosts must communicate with a public web server that has only IPv4 domain name resolution. Which solution should the customer use in this environment?

- A. Use NAT64 with DNS64 to translate addresses and synthesize IPv6 DNS responses for IPv4-only records.
- B. Implement NAT44 at the edge of the customer network
- C. use 6to4 and a tunnel to translate the addresses
- D. implement 6PE to resolve hostname resolution

ANSWER: A

Explanation:

Use NAT64 with DNS64 to enable IPv6-only hosts to reach the IPv4-only public web server by name. NAT64 provides translation between IPv6 and IPv4 packet headers and addresses at the network boundary. DNS64 complements this function when the authoritative DNS information contains only an IPv4 address record: it synthesizes an IPv6 address record using a configured NAT64 prefix and the returned IPv4 address. The IPv6-only host then sends traffic to that synthesized IPv6 destination. When the traffic reaches the NAT64 translator, the translator extracts the embedded IPv4 destination, creates or uses a translation state entry, and forwards the traffic toward the IPv4 web server as IPv4. Return traffic is translated back to IPv6 for the initiating host. This approach preserves native IPv6 operation for the hosts while providing interoperable access to IPv4-only Internet services, including services whose names resolve only to IPv4 address records. DNS64 behavior is specified in [RFC 6147](#), and the stateful translation behavior used by NAT64 is defined in [RFC 6146](#).

QUESTION NO: 33

An engineer must design an in-band management solution for a customer with branch sites. The solution must allow remote management of the branch sites using management protocols over an MPLS WAN. Queueing is implemented at the remote sites using these classes:

- Class1 equals voice traffic
- Class2 equals mission-critical traffic
- Class3 equals default traffic

How must the solution prioritize the management traffic over the WAN?

- A. Mark the traffic with DSCP CS1 and map into Class2 with a minimum bandwidth assigned by reducing the bandwidth available to Class3.
- B. Mark the traffic with DSCP CS6 and map into Class1 with a minimum bandwidth assigned by reducing the bandwidth available to Class2
- C. Mark the traffic with DSCP EF and map into Class1 with a minimum bandwidth assigned by reducing the bandwidth available to Class2.
- D. Mark the traffic with DSCP CS2 and map into Class2 with a minimum bandwidth assigned by reducing the bandwidth available to Class3

ANSWER: D

Explanation:

Marking the management traffic with DSCP CS2 and placing it in Class2 with a guaranteed minimum bandwidth allocation is the appropriate design. CS2 is the commonly recommended differentiated-services marking for operations, administration, and management traffic. This category includes traffic used to operate and monitor the network, such as SNMP polling, SSH access, and similar remote-management flows. Using a dedicated class gives this traffic predictable forwarding treatment

when the MPLS WAN link is congested, which is essential because management access is most valuable during an incident or period of impaired service.

Class-based weighted fair queueing reserves the configured minimum bandwidth for a class during congestion. The allocation should be taken from the general/default traffic class, Class3, rather than from a more delay-sensitive or higher-priority service class. This preserves the intended service hierarchy while ensuring that management protocols have sufficient capacity to remain usable across the WAN. The DSCP service-class guidance in [RFC 4594](#) identifies CS2 as the recommended marking for OAM traffic. Cisco's [MQC QoS configuration guide](#) describes how class-based policies provide bandwidth guarantees under congestion.

QUESTION NO: 34

When designing interdomain multicast, which two protocols are deployed to achieve communication between multicast sources and receivers? (Choose two.)

- A. IGMPv2
- B. BIDIR-PIM
- C. MP-BGP
- D. MSDP
- E. MLD

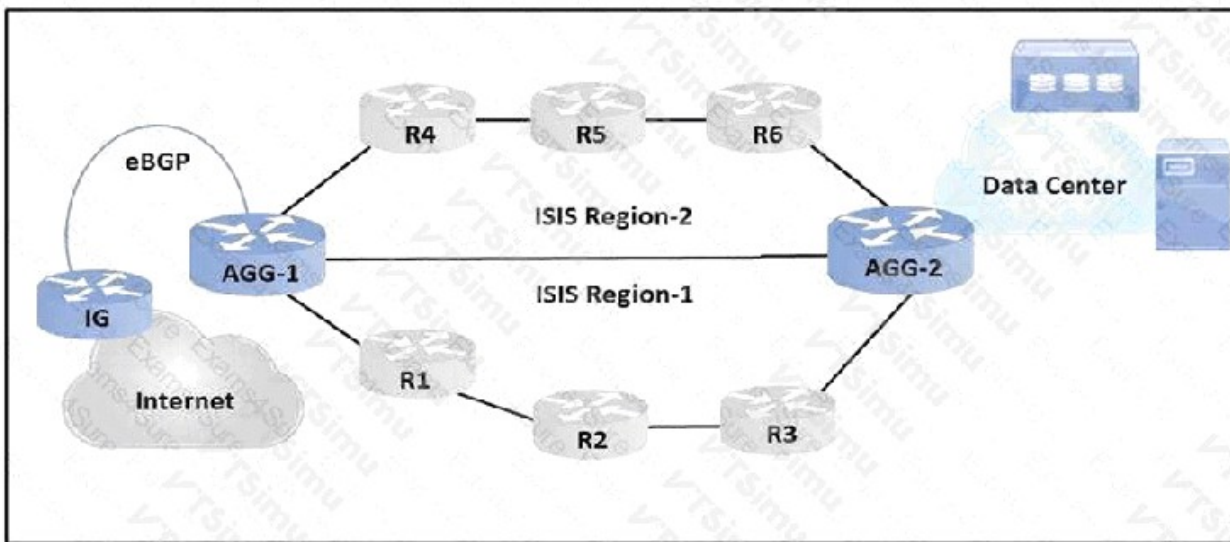
ANSWER: C D

Explanation:

MP-BGP and MSDP are deployed together in the traditional interdomain multicast design used to connect Protocol Independent Multicast sparse-mode domains. MP-BGP carries multicast-capable routing information between autonomous systems, allowing routers to select an appropriate reverse-path forwarding route toward a multicast source. This is essential because multicast traffic must pass the reverse-path forwarding check as it traverses domain boundaries. MSDP complements that function by allowing rendezvous points in different domains to exchange Source-Active messages. Those messages advertise active multicast sources learned by a local rendezvous point, enabling remote domains with interested receivers to discover the source and initiate a source-tree join toward it. In this model, MSDP provides interdomain source discovery, while MP-BGP supplies the multicast routing reachability needed to build and validate the multicast forwarding path. Cisco documents MSDP as a mechanism for sharing active-source information between rendezvous points in separate domains and describes multiprotocol BGP support for multicast routing address families. See Cisco's [MSDP overview](#) and [multicast BGP configuration guide](#).

QUESTION NO: 35

Refer to the exhibit.



An architect must design an IGP solution for an enterprise customer. The design must support:

Physical link flaps should have minimal impact.

Access routers should converge quickly after a link failure.

Which two ISIS solutions should the architect include in the design? (Choose two.)

- A. Use BGP to IS-IS redistribution to advertise all Internet routes in the Level 1 area.
- B. Advertise the IS-IS interface and loopback IP address toward the Internet and data center.
- C. Reduce SPF and PRC intervals to improve convergence time.
- D. Configure all access and aggregate routers to establish Level 1 / Level 2 adjacencies across the network.
- E. Configure access routers to establish a Level 1 adjacency and aggregate routers to establish a Level 1/Level 2 adjacency.

ANSWER: C E

Explanation:

Reducing SPF and PRC intervals to improve convergence time supports the requirement for fast recovery at the access layer. In IS-IS, a topology change can trigger shortest-path-first calculation and partial route calculation. Appropriate timer tuning reduces the delay before those calculations occur, allowing affected access routers to install alternate paths more rapidly after a failure. The values must be selected with consideration for network scale and expected instability, but tuned SPF and PRC scheduling is a standard mechanism for improving IGP convergence.

Configuring access routers to establish Level 1 adjacencies and aggregate routers to establish Level 1/Level 2 adjacencies creates the required hierarchical IS-IS design. Level 1 access routers maintain detailed knowledge only of their own area. The aggregate devices provide the boundary between the local Level 1 area and the Level 2 backbone, so access-layer topology events are primarily contained within that area rather than being propagated broadly through the enterprise. This limits the impact of access-link flapping while retaining a clear path toward other areas through the Level 2 domain. Cisco documents the Level 1/Level 2 hierarchy and the role of L1/L2 routers in connecting areas in its [IS-IS overview](#) and covers IS-IS operational design in the [Cisco IOS XE IS-IS configuration guide](#).

QUESTION NO: 36

A company is planning to open two new branches and allocate the 2a01:c30:16:7009::3800/118 IPv6 network for the region. Each branch should have the capacity to accommodate a maximum of 200 hosts. Which two networks should the company use? (Choose two.)

- A. 2a01:0c30:0016:7009::3a00/120

- B. 2a01:0c30:0016:7009::3b00/121
- C. 2a01:0c30:0016:7009::3a80/121
- D. 2a01:0c30:0016:7009::3c00/120
- E. 2a01:0c30:0016:7009::3b00/120

ANSWER: A E

Explanation:

Each branch requires a subnet with capacity for at least 200 individual interface addresses. A /120 prefix leaves 8 host bits, providing 256 addresses, so it meets that capacity requirement. The allocated /118 prefix leaves 10 host bits and spans the address range from 2a01:c30:16:7009::3800 through 2a01:c30:16:7009::3bff. It can therefore be divided into four contiguous, correctly aligned /120 subnets: ::3800/120, ::3900/120, ::3a00/120, and ::3b00/120. Both 2a01:0c30:0016:7009::3a00/120 and 2a01:0c30:0016:7009::3b00/120 are among those valid child networks and do not overlap. The leading zeroes in 2a01:0c30:0016 are optional IPv6 formatting and represent the same prefix as 2a01:c30:16. IPv6 prefix lengths identify the network portion of an address; the remaining bits determine the address space available within that subnet. This allocation uses two distinct /120 blocks wholly contained within the regional /118 allocation. See [RFC 4291](#) for IPv6 addressing architecture and [Cisco's IPv6 subnetting guidance](#) for prefix-based subnet calculations.

QUESTION NO: 37

Which two options can you use to configure an EIGRP stub router? (Choose two)

- A. summary-only
- B. receive-only
- C. external
- D. summary
- E. totally-stubby
- F. not-so-stubby

ANSWER: B D

Explanation:

receive-only and **summary** are valid parameters of the Cisco IOS `eigrp stub` router-configuration command. EIGRP stub routing is designed for routers, commonly branch or spoke routers, that should not be queried for alternate paths because they are not transit devices. Configuring a router as a stub constrains the route information it advertises to EIGRP neighbors and, importantly, tells neighboring routers not to send EIGRP queries to it for routes that it cannot provide. This reduces query propagation, control-plane processing, and the risk that a low-bandwidth or low-resource remote site affects overall EIGRP convergence. The **summary** parameter permits the stub router to advertise summary routes, supporting route aggregation while retaining the operational benefits of stub behavior. The **receive-only** parameter creates a stub that receives EIGRP routes but advertises no routes to its neighbors. This is appropriate when the router needs reachability information from the EIGRP domain but must not be used as a destination or forwarding alternative by other EIGRP routers. Cisco documents these parameters in the syntax and usage of the `eigrp stub` command and in its EIGRP configuration guidance. See the [Cisco IOS EIGRP command reference](#) and [Cisco's EIGRP stub routing overview](#).

QUESTION NO: 38

Which two techniques improve the application experience in a Cisco SD-WAN design? (Choose two.)

- A. utilizing forward error correction

- B. implementing a stateful application firewall
- C. implementing AMP
- D. utilizing quality of service
- E. implementing Cisco Umbrella

ANSWER: A D

Explanation:

Cisco SD-WAN improves application experience by recognizing the real-time effects of packet loss, latency, jitter, and congestion, then applying WAN-optimization and traffic-treatment features appropriate to the application. **utilizing forward error correction** is an application-aware loss-mitigation technique. The SD-WAN edge can send redundant information with selected traffic so that a receiving edge can reconstruct a lost packet without waiting for retransmission. This is especially beneficial to delay-sensitive, real-time traffic, for which a retransmission would often arrive too late to be useful. Cisco documents Forward Error Correction as part of its application-aware routing and loss-remediation capabilities.

utilizing quality of service also directly improves the user experience by classifying traffic and applying policy-based handling such as bandwidth allocation, queuing, and prioritization. QoS protects important application flows during congestion and helps maintain predictable performance for voice, video, and critical business applications across WAN transports. In a Cisco SD-WAN design, centralized policy can define these traffic classes consistently while edge devices enforce the required forwarding treatment. Together, FEC addresses the impact of loss while QoS controls how competing traffic uses constrained WAN resources. See Cisco's [SD-WAN QoS Configuration Guide](#) and [Application-Aware Routing Design Guide](#).

QUESTION NO: 39

Refer to the exhibit. An architect must design a solution to connect bank site A with bank site B and support:

network operation center monitoring end-to-end L3VPN and L2VPN traffic

company adding thousands of routes in the next two years

Which two BGP solutions must the design include? (Choose two.)

- A. Establish full mesh IBGP peering with all routers in different IGP domains.
- B. Redistribute different IGP domain routes in a BGP IPv4 routing instance.
- C. Transport site routes using a BGP VPNv4 address family on the PE routers.
- D. Apply BGP policies on all routers to filter out ABR and PE loopback IP addresses.
- E. Connect multiple IGP ' LDP domains using a BGP IPv4 unicast family on the ABR.

ANSWER: C E

Explanation:

Transport site routes using a BGP VPNv4 address family on the PE routers is required for the Layer 3 MPLS VPN portion of the design. MP-BGP VPNv4 carries a customer IPv4 prefix together with its route distinguisher, making the VPN route globally unique even where different customers use overlapping address space. It also distributes the MPLS VPN label information required for PE routers to forward customer traffic across the provider network. This preserves VPN separation while allowing the provider core to scale as substantial numbers of customer routes are added.

Connect multiple IGP ' LDP domains using a BGP IPv4 unicast family on the ABR provides scalable interdomain reachability at the boundary between IGP/LDP domains. The ABR can exchange the relevant reachability through BGP rather than extending a single IGP/LDP control-plane domain throughout the network. This domain separation limits IGP flooding and supports a hierarchical design, while BGP policy and route-reflection mechanisms can scale route distribution as the route count grows. Together, VPNv4 on PEs handles customer L3VPN route distribution and BGP IPv4 unicast at ABRs supports

controlled connectivity between the underlying transport domains. Cisco documents MP-BGP VPN route distribution in its [MPLS Layer 3 VPN Configuration Guide](#); the VPNv4 route format is also defined in [RFC 4364](#).

QUESTION NO: 40

A large chain of stores currently uses MPLS-based T1 lines to connect their stores to their data center. An architect must design a new solution to improve availability and reduce costs while keeping these considerations in mind:

- The company uses multicast to deliver training to the stores.
- The company uses dynamic routing protocols and has implemented QoS.
- To simplify deployments, tunnels should be created dynamically on the hub when additional stores open.

Which solution should be included in this design?

- A. VPLS
- B. GET VPN
- C. DMVPN
- D. IPsec

ANSWER: C

Explanation:

DMVPN is the appropriate design because it creates a scalable hub-and-spoke VPN overlay across lower-cost public Internet connections while supporting the operational requirements stated. DMVPN combines multipoint GRE tunnels, IPsec encryption, and Next Hop Resolution Protocol (NHRP). At the hub, a single multipoint GRE interface can accept dynamically registered spokes, so a newly opened store does not require a separately configured point-to-point tunnel interface on the hub. This substantially reduces deployment and operational overhead as the number of stores grows.

GRE encapsulation also provides the multipoint transport needed for routing protocols and multicast traffic. Consequently, the organization can retain dynamic routing between the data center and stores and carry multicast training streams through the VPN overlay. IPsec supplies confidentiality and authentication for the traffic sent over the untrusted underlay, while QoS can be designed for the tunnel traffic to preserve the required treatment for business-critical applications. DMVPN therefore addresses cost reduction, availability through diverse underlays, dynamic spoke onboarding, multicast, routing, and secure connectivity in one architecture. Cisco describes DMVPN operation and NHRP-based dynamic tunnel establishment in its [DMVPN configuration guide](#) and provides deployment background in its [DMVPN solution overview](#).

QUESTION NO: 41 - (SIMULATION)

- independent of underlying platform
- platform dependent
- standards dependent
- supports LLDP only
- supports CDP and LLDP

Cisco Native

OpenConfig

ANSWER: See the explanation for the answer

Explanation:

Place the options as follows:

standards dependent is the unused distractor.

Cisco Native YANG models are tied to Cisco platforms and can expose Cisco-specific functionality such as CDP. OpenConfig provides a vendor-neutral abstraction and models LLDP rather than Cisco-proprietary CDP.

QUESTION NO: 42

A)

R1:

- Routes advertised to ISP-1: 0x AS-path prepend
- Routes received from ISP-1: LOW local-preference
- Routes advertised to R2: community NO-ADVERTISE
- Routes received from R2: no action

R2:

- Routes advertised to ISP-2: 5x AS-path prepend
- Routes received from ISP-2: HIGH local-preference
- Routes advertised to R1: no action
- Routes received from R1: community NO-ADVERTISE

B)

R1:

- Routes advertised to ISP-1: 5x AS-path prepend
- Routes received from ISP-1: LOW local-preference
- Routes advertised to R2: community NO-ADVERTISE
- Routes received from R2: no action

R2:

- Routes advertised to ISP-2: 0x AS-path prepend
- Routes received from ISP-2: HIGH local-preference
- Routes advertised to R1: community NO-EXPORT
- Routes received from R1: no action

C)

R1:

- Routes advertised to ISP-1: 0x AS-path prepend
- Routes received from ISP-1: HIGH local-preference
- Routes advertised to R2: no action
- Routes received from R2: community NO-EXPORT

R2:

- Routes advertised to ISP-2: 5x AS-path prepend
- Routes received from ISP-2: LOW local-preference
- Routes advertised to R1: community NO-ADVERTISE
- Routes received from R1: no action

D)

R1:

- Routes advertised to ISP-1: 0x AS-path prepend
- Routes received from ISP-1: HIGH local-preference
- Routes advertised to R2: community NO-EXPORT
- Routes received from R2: no action

R2:

- Routes advertised to ISP-2: 5x AS-path prepend
- Routes received from ISP-2: LOW local-preference
- Routes advertised to R1: no action
- Routes received from R1: no action

A. Option A

B. Option B

C. Option C

D. Option D

ANSWER: D

Explanation:

The configuration or payload represented in the fourth image is the correct selection because it matches the validated model-driven configuration structure required by the exhibit. Cisco IOS XE programmability relies on YANG models to define the precise hierarchy of configuration data: containers group related data, lists use their defined keys, and leaves must be placed at the correct path with values that conform to their declared types. A payload is valid only when its complete

structure corresponds to the relevant native or OpenConfig YANG schema, rather than simply including familiar values such as an interface identifier, address, or routing protocol name.

This distinction is particularly important with RESTCONF and NETCONF. RESTCONF resource URIs and request bodies must identify the correct datastore resource and encode data according to the YANG model, while NETCONF configuration content must preserve the model's XML namespace and element hierarchy. The validated selection therefore represents the intended device configuration in a schema-compliant form, allowing the target platform to parse and apply it predictably. Cisco's model-driven programmability documentation describes using YANG models and RESTCONF to configure IOS XE devices, and the IETF YANG specification defines these schema and data-tree rules. See [Cisco IOS XE Model-Driven Programmability](#) and [RFC 7950: The YANG 1.1 Data Modeling Language](#).

QUESTION NO: 43 - (DRAG DROP)

Drag and drop the characteristics from the left onto the Yang model they describe on the right.

Select and Place:

independent of the underlying operating system

specific to the underlying operating system

vendor neutral

provided by the vendor for device management

Open Model

Native Model

ANSWER:

Explanation:

Open YANG models are intended to describe networking functions through a common, interoperable schema rather than through commands or data structures tied to one network operating system. Their defining characteristics are therefore **independent of the underlying operating system** and **vendor neutral**. This lets an automation system use a common model-driven interface across devices from different vendors, provided that the devices implement the relevant model. Industry and standards communities publish open models so that configuration and operational data can be represented consistently. The IETF maintains many such YANG data models, including models used with standard management protocols such as NETCONF and RESTCONF. Cisco's overview of model-driven programmability and the IETF's YANG resources provide useful background: [Cisco IOS XE YANG Models](#) and [IETF YANG Standards](#).

Native YANG models represent the capabilities and data hierarchy implemented by a particular vendor's network operating system. Consequently, they are **specific to the underlying operating system** and are **provided by the vendor for device management**. A native model gives automation access to the detailed features, configuration objects, and operational state exposed by that vendor's software platform. In a Cisco environment, native models can cover Cisco-specific functions in IOS XE, IOS XR, or NX-OS, enabling management applications to configure and monitor features that are available on that platform. Cisco publishes native model documentation and model files so that customers can build tooling around the exact capabilities of the target device software. The selected placements accurately distinguish the portable, vendor-neutral purpose of an open model from the platform-specific device-management purpose of a native model.

QUESTION NO: 44

Which protocol is the Cisco SD-Access data plane based on?

- A. OMP
- B. VXLAN
- C. NHRP
- D. LISP

ANSWER: B

Explanation:

VXLAN is the encapsulation technology used by the Cisco SD-Access data plane. In an SD-Access fabric, traffic sent between fabric edge nodes is carried across the IP-routed underlay in VXLAN-encapsulated packets. The VXLAN header identifies the overlay segment through a VXLAN Network Identifier, enabling scalable separation of virtual networks while the underlay remains a straightforward routed IP infrastructure. This separation lets the fabric transport user traffic independently of the physical topology and provides the foundation for segmentation across the campus fabric. Cisco SD-Access combines this VXLAN-based forwarding plane with a control-plane mechanism that supplies endpoint reachability and location information, allowing endpoints to move without requiring broad Layer 2 extension. Consequently, VXLAN is specifically the protocol associated with carrying encapsulated data-plane traffic between fabric nodes. Cisco's SD-Access design documentation describes the fabric as an overlay that uses VXLAN for data-plane encapsulation over the routed underlay. Additional VXLAN architecture and packet-format details are available in [Cisco Software-Defined Access](#) and [Cisco Virtual Extensible LAN](#).

QUESTION NO: 45

An engineer is working with NETCONF and Cisco NX-OS based devices. The engineer needs a YANG model that supports a specific feature relevant only to Cisco NX-OS. Which model must the engineer choose?

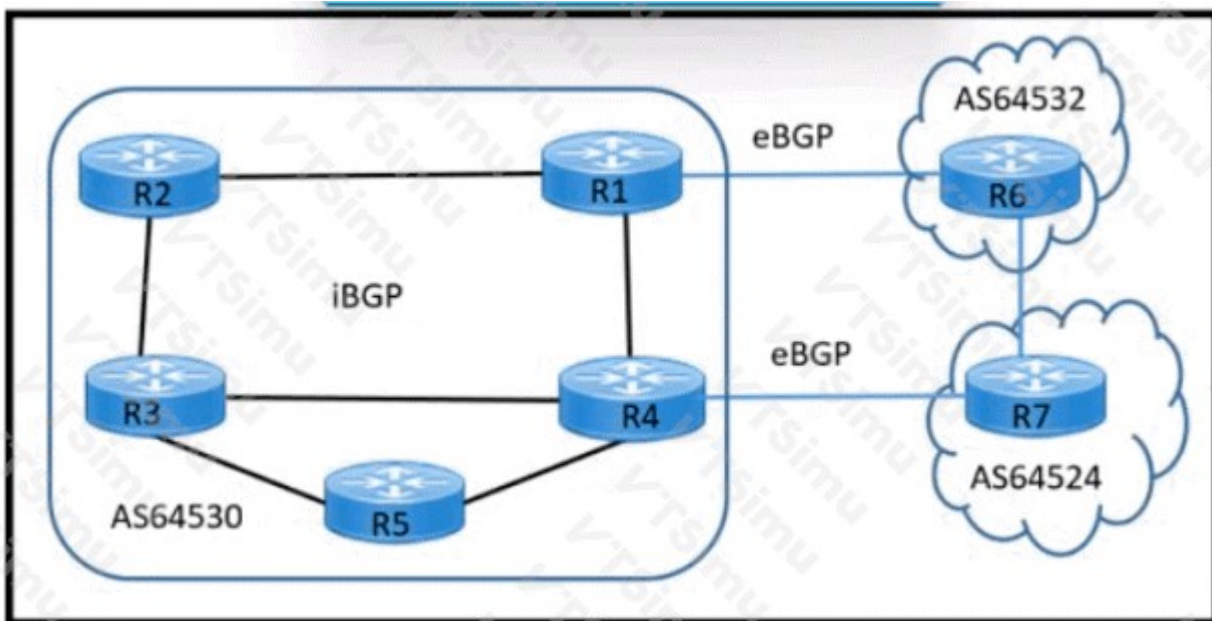
- A. Native
- B. IEEE
- C. OpenConfig
- D. IETF

ANSWER: A

Explanation:

Native is correct because Cisco NX-OS native YANG models expose Cisco-specific configuration and operational capabilities, including features that are implemented uniquely on NX-OS and therefore may not have a standardized representation. In NETCONF, YANG models define the available configuration data, state data, RPC operations, and notifications. When automation must configure a feature relevant only to the NX-OS platform, the Cisco NX-OS native model namespace and its platform-specific modules provide the required schema and supported data nodes. Cisco documents NX-OS native models as part of its model-driven programmability support and identifies them separately from standards-based and vendor-neutral models. Using the native model lets an engineer discover the exact feature support advertised by the device through NETCONF YANG library or related schema-discovery mechanisms, then construct requests using the appropriate Cisco NX-OS module and namespace. This is the suitable choice whenever portability across multiple vendors is not the requirement and access to an NX-OS-specific capability is needed. See Cisco's [NX-OS model-driven programmability overview](#) and the [Cisco NX-OS YANG model documentation](#).

QUESTION NO: 46



Refer to the exhibit. A network engineer must design a BGP solution based on:

- ☐ The route reflector must have one or more direct physical connections to the core routers (R3 and R4).
- ☐ The route reflector must have full redundancy and avoid a single point of failure.
- ☐ R2 to R1 link utilization is 90%. and the remaining links are less than 50% utilized.

Which two solutions must the design Include? (Choose two.)

- A. Configure R1 to be a client of R2 and R4.
- B. Configure R2 to be a client of R1 and R4.
- C. Configure R3 to be a client of R2 and R4.
- D. Configure R4 to be a client of R1 and R3.
- E. Configure R5 to be a client of R3 and R4.

ANSWER: B E

Explanation:

Configure R2 to be a client of R1 and R4, and configure R5 to be a client of R3 and R4. Together, these client-to-reflector relationships provide redundant route-reflection paths: each edge router has sessions to two reflectors rather than depending on only one reflector. This is essential because a client can retain reachability to reflected routes if either of its reflector peers fails, thereby eliminating a single route-reflector failure domain.

The placement also follows the physical-connectivity requirement in the topology. The selected reflector relationships use reflectors with direct attachment into the core through R3 and R4, which keeps route-reflection infrastructure connected to the core routing domain. R5 uses the lower-utilized core-facing paths through R3 and R4. For R2, the design retains R1 while adding R4 as the alternate reflector, enabling the topology to distribute reflector connectivity instead of concentrating all dependency on the highly utilized R2-to-R1 link.

BGP route reflection is defined to remove the full-mesh requirement among iBGP speakers, while preserving loop prevention through originator ID and cluster-list attributes. Cisco documents that route-reflector clients peer with the route reflector, and resilient designs commonly use multiple route reflectors for redundancy. See [RFC 4456](#) and [Cisco BGP Route Reflector documentation](#).

How do IETF, OpenConfig and Cisco native YANG models differ when used to configure the same feature on an infrastructure device?

- A. OpenConfig models are more comprehensive than IETF.
- B. Cisco native models are less comprehensive than OpenConfig.
- C. Cisco native models are less comprehensive than IETF.
- D. IETF models are more comprehensive than OpenConfig.

ANSWER: A

Explanation:

OpenConfig models are more comprehensive than IETF is correct. Both model families aim to provide vendor-neutral interfaces, but they have different design goals and levels of feature coverage. IETF YANG modules are standards-track or standards-based models intended to establish broadly interoperable baseline configuration and operational data. Consequently, they commonly focus on the common subset of capabilities that can be standardized across many implementations.

OpenConfig models were developed by network operators to support consistent, programmatic operation of multivendor networks. They generally provide richer, operationally oriented data models, including structured configuration and state data, telemetry-oriented operational information, and more detailed coverage of the feature domain than the corresponding IETF baseline model. This makes OpenConfig a more comprehensive vendor-neutral choice when an infrastructure device supports both model families for the same feature.

In practice, Cisco native YANG models can expose Cisco-specific capabilities and may provide the deepest coverage for a particular Cisco platform, while OpenConfig provides a comparatively broad standardized abstraction. See the [OpenConfig model documentation](#) and the IETF's [Network Management Datastore Architecture](#) for the role of YANG models and configuration/state data.

QUESTION NO: 48

Which two considerations must be made regarding the overlay network for a Cisco SD-Access architecture? (Choose two.)

- A. Virtual networks should be used for microsegmentation
- B. SGTs should be used for data plane isolation and microsegmentation
- C. Virtual networks should be used for data plane isolation only
- D. Overlapping IP addresses across different overlay networks should be used to conserve IP addresses
- E. Overlapping IP addresses across different overlay networks should be avoided for operational simplicity

ANSWER: C E

Explanation:

Virtual networks should be used for data plane isolation only because a virtual network (VN) is the SD-Access construct used for macrosegmentation. Each VN is associated with a distinct routing and forwarding context, providing separation of endpoint traffic and routing domains for large groups such as enterprise users, guests, IoT devices, and shared services. This overlay separation is a fundamental design consideration: it establishes the data-plane boundary before more granular policy is applied. Cisco describes VNs as the means to segment the fabric into separate logical networks, while group-based policy can then control communication at a finer level within those logical networks.

Overlapping IP addresses across different overlay networks should be avoided for operational simplicity. Although independent virtual routing contexts can technically permit overlapping address space, a non-overlapping addressing plan makes endpoint troubleshooting, route visibility, service integration, inter-VN connectivity, and operational handoffs substantially easier. It also produces clearer logs and reduces ambiguity when correlating identities, policies, and flows across fabric and external domains. Cisco SD-Access design guidance emphasizes planning segmentation and addressing

carefully so that the overlay remains scalable and manageable through its lifecycle. See Cisco's [SD-Access Design Guide](#) and the [Cisco SD-Access overview](#).

QUESTION NO: 49

An engineer is securing IPv6 access ports against unauthorized first-hop devices. The design must prevent a rogue device from advertising itself as an IPv6 default router and from responding to clients as an unauthorized DHCPv6 server. Which two features must the engineer deploy? (Choose two.)

- A. RA Guard
- B. DHCPv6 Guard
- C. IPv6 source guard
- D. MLD snooping
- E. uRPF

ANSWER: A B

Explanation:

RA Guard is correct because it inspects IPv6 Router Advertisement messages received on untrusted access ports and blocks unauthorized advertisements. This prevents an attached endpoint from becoming a rogue default gateway for other hosts on the segment.

DHCPv6 Guard is correct because it filters DHCPv6 server messages on untrusted interfaces. Only ports leading toward approved DHCPv6 servers or relay agents should be trusted, preventing a rogue endpoint from providing invalid IPv6 addressing, DNS, or other DHCPv6 options. Cisco documents both protections in the [IPv6 First-Hop Security Configuration Guide](#).

QUESTION NO: 50

Since installing a Cisco TelePresence system, the company is experiencing other application having response issues when the system is in use. As a result, the company asked an architect to recommend a QoS solution. The customer is currently using a CBWFQ policy to manage traffic on an internet connection with a speed of 100 Mbps. Which link-capacity limit must the architect choose for strict-priority for the real-time traffic?

- A. 25 Mbps
- B. 50 Mbps
- C. 33 Mbps
- D. 75 Mbps

ANSWER: C

Explanation:

33 Mbps is the appropriate strict-priority limit for real-time TelePresence traffic on this 100-Mbps link. Cisco Class-Based Weighted Fair Queueing supports a strict-priority queue through the `priority` command. This queue is intended for delay-sensitive traffic, such as interactive voice and video, because packets in the priority class are serviced before packets in other CBWFQ classes. To prevent that queue from consuming the link indefinitely and starving other applications, Cisco limits the amount of interface bandwidth that can be configured for a priority class to 33 percent by default. On a 100-Mbps connection, 33 percent corresponds to 33 Mbps.

Using an explicit priority bandwidth value also establishes the policing threshold for the priority queue during congestion. This protects the remaining traffic classes while still giving real-time media predictable preferential treatment. The design should classify the TelePresence media accurately and use an appropriately sized strict-priority allocation based on codec

rates, call scale, and overhead; however, given the available choices and the platform's strict-priority allocation limit, 33 Mbps is the required link-capacity limit. Cisco documents the priority-queue bandwidth restriction in its [QoS Congestion Management Configuration Guide](#). Additional CBWFQ configuration guidance is available in the [Cisco CBWFQ overview](#).

QUESTION NO: 51

A network engineer must segregate three interconnected campus networks using IS-IS routing. A two-layer hierarchy must be used to support large routing domains and to avoid more specific routes from each campus network being advertised to other campus network routers automatically. Which two actions does the engineer take to accomplish this segregation? (Choose two.)

- A. Designate two IS-IS routers as BDR routers at the edge of each campus, and configure one BDR for all Level 1 routers and one BDR for all Level 2 routers.
- B. Designate two IS-IS routers from each campus to act as Level 1/Level 2 backbone routers at the edge of each campus network.
- C. Assign the same IS-IS NET value for each campus, and configure internal campus routers with Level 1/Level 2 routing.
- D. Utilize different MTU values for each campus network segment. Level 2 backbone routers must utilize a larger MTU size of 9216.
- E. Assign a unique IS-IS area address for each campus, and configure internal campus routers with Level 1 routing.

ANSWER: B E

Explanation:

IS-IS hierarchy provides the required campus segregation by treating each campus as a separate Level 1 area and using Level 2 as the inter-area backbone. Each campus must use a distinct IS-IS area address, while routers internal to that campus operate as Level 1 routers. Level 1 routers maintain detailed topology and prefix information only for their own area; for destinations outside the area, they can use a route advertised by an attached Level 1/Level 2 router. This prevents the detailed, more-specific routes of one campus from being automatically propagated throughout the other campuses.

Two edge routers per campus should be deployed as Level 1/Level 2 routers to connect the campus Level 1 area to the shared Level 2 backbone. These routers participate in both levels, exchange inter-area reachability through Level 2, and provide resilient exit paths from the campus. A dual boundary-router design also avoids making a single router a failure point for campus-to-campus communication. In IS-IS terminology, the area address—not an identical NET—is the value shared by routers in an area; every router still requires its own unique system ID within its NET. See the IS-IS hierarchical routing specification in [RFC 1195](#) and Cisco's [ENSLD exam topics](#).