

Implementing Cisco Service Provider Advanced Routing Solutions

Cisco 300-510

Version Demo

Total Demo Questions: 40

Total Premium Questions: 409

Buy Premium PDF

<https://passqueen.com>

support@passqueen.com

Topic Break Down

Topic	No. of Questions
Topic 1, Routing Protocols	263
Topic 2, MPLS	57
Topic 3, Segment Routing	57
Topic 4, Services	31
Topic 5, Mix Questions	1
Total	409

QUESTION NO: 1

While configuring Cisco NSF awareness, a network engineer enters the `bgp graceful-restart` command after the BGP session is established in a router that runs IOS XE Software. Graceful restart capabilities are not exchanged. Which two actions should be taken? (Choose two.)

- A. Reload the router
- B. Verify that BGP route dampening is configured
- C. Reduce BGP convergence time
- D. Issue the `clear ip bgp *` command
- E. Issue the `show ip bgp neighbors` command.

ANSWER: A D

Explanation:

BGP graceful restart capability is negotiated as part of the BGP OPEN-message capability exchange, which occurs when a BGP peering session is established. Configuring `bgp graceful-restart` after the session is already up does not retroactively send the graceful-restart capability to the neighbor. The session must therefore be reestablished so that both peers can exchange and record their graceful-restart capabilities during a new OPEN-message exchange.

Reload the router is a valid way to force the established BGP sessions to restart. When the router returns to service, BGP peerings are established again and the configured graceful-restart capability can be advertised. Issuing `clear ip bgp *` is also valid because it clears the BGP peer sessions, causing them to reconnect and renegotiate capabilities without requiring a full device reload. Cisco specifically documents reloading the router or clearing the BGP session when graceful restart is configured after a BGP session has already been established. See Cisco's [BGP Nonstop Forwarding Awareness documentation](#) and the [BGP Graceful Restart RFC](#).

QUESTION NO: 2

Refer to the exhibit.

Which statement about this configuration is true?

- A. Router 1 sends and receives multiple best paths from neighbor 192.168.1.1
- B. Router 1 sends up to two paths to neighbor 192.168.1.1 for all routes
- C. Router 1 receives up to two paths from neighbor 192.168.1.1 for all routes in the same AS
- D. Router 1 receives only the best path from neighbor 192.168.1.1

ANSWER: A

Explanation:

Router 1 sends and receives multiple best paths from neighbor 192.168.1.1 because BGP Additional Paths is configured to exchange more than the conventional single best path between BGP peers. BGP normally advertises only its selected best path for a destination, even when it has learned multiple eligible paths. With Additional Paths, a router can select multiple paths for advertisement and can accept multiple paths received from a peer. The paths are distinguished by path identifiers, allowing Router 1 to retain multiple routes for the same prefix from 192.168.1.1 rather than replacing all but one in the BGP table.

The configured best-path selection limit controls how many eligible best paths Router 1 can advertise, while the neighbor Additional Paths send and receive capabilities enable the bidirectional exchange. This is useful in service-provider designs such as route reflection, where advertising more than one path improves path visibility and can reduce path hiding. Both peers must negotiate the Add-Path capability during BGP session establishment for the additional path identifiers and routes to be exchanged. Cisco documents the configuration and operational behavior in its [BGP Additional Paths Configuration Guide](#); the protocol behavior is standardized in [RFC 7911](#).

QUESTION NO: 3 - (DRAG DROP)

DRAG DROP

Drag and drop the attributes for the BGP route selection on the left into the correct order on the right. Not all options are used.

Select and Place:

lowest router ID	Step 1
highest router ID	Step 2
lowest local preference	Step 3
if both paths are external, prefer the newest path	Step 4
highest local preference	Step 5
highest weight	Step 6
eBGP over iBGP paths	
if both paths are external, prefer the oldest path	
lowest MED	

ANSWER:



Explanation:

Cisco BGP applies its best-path selection criteria in a fixed order, continuing only while multiple otherwise valid paths remain tied. From the attributes supplied here, the first applicable criterion is highest weight. Weight is a Cisco-specific, local value and a larger value is preferred. The next applicable criterion is highest local preference. Local preference is propagated within an autonomous system to express the preferred outbound exit path, so the route carrying the larger local-preference value is selected.

When those attributes do not separate the paths, BGP compares MED and selects the path with the lowest MED. MED is normally used as an indication of a preferred entry point into a neighboring autonomous system. If a tie still exists, an external BGP path is preferred over an internal BGP path. Where eligible external paths still tie, Cisco can prefer the oldest path, which helps preserve route stability by retaining an established external path rather than unnecessarily changing to another equivalent one.

The lowest router ID is a later deterministic tie-breaker. It is used only after the preceding relevant attributes have failed to choose a single best path, ensuring consistent selection when otherwise equivalent candidate paths remain. This produces the required order: highest weight, highest local preference, lowest MED, eBGP over iBGP paths, if both paths are external prefer the oldest path, and lowest router ID. Cisco documents the BGP best-path algorithm and its ordered comparisons in its [BGP Best Path Selection Algorithm](#) technical guide.

QUESTION NO: 4

What are two differences between OSPF and IS-IS? (Choose two.)

- A. Unlike IS-IS routers, an OSPF router belongs to only one area in addition to the backbone area
- B. OSPF uses a router ID to identify a router, and IS-IS uses a system ID

C. OSPF is a link-state routing protocol, and IS-IS is a distance-vector routing protocol

D. Unlike OSPF, IS-IS supports virtual links

E. OSPF elects a DR and BDR, and IS-IS elects a DIS

ANSWER: B E

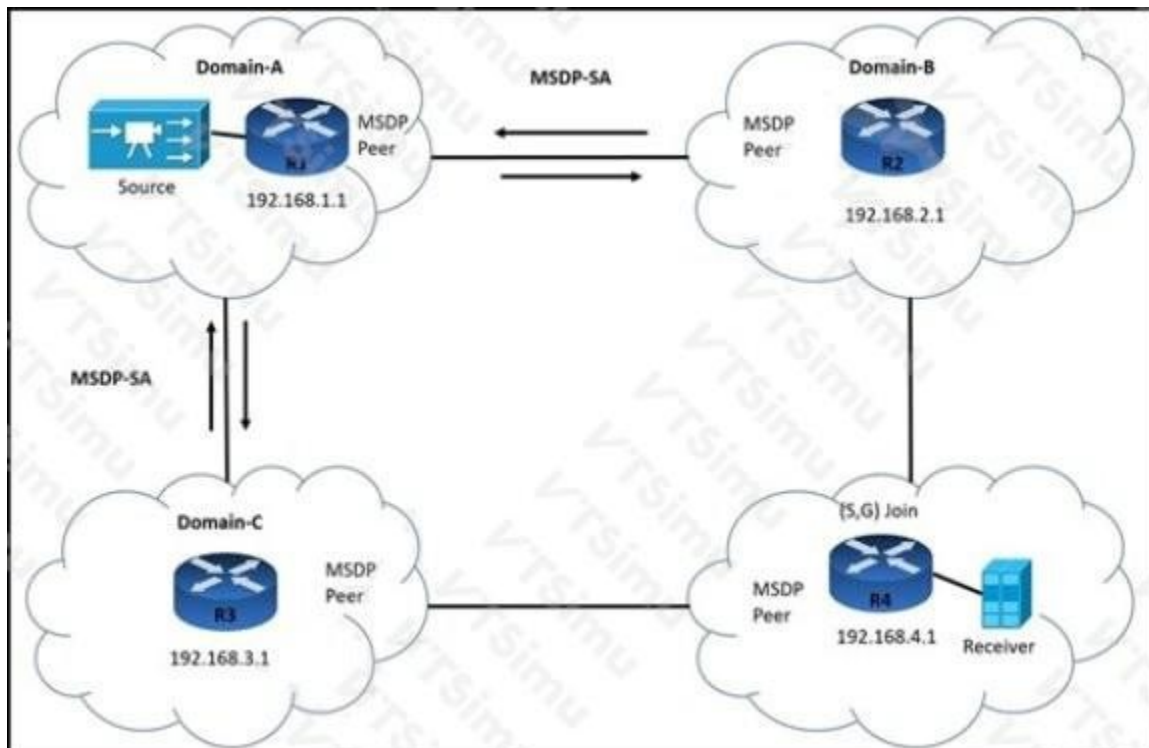
Explanation:

OSPF uses a router ID to identify a router, and IS-IS uses a system ID. In OSPF, the Router ID is a 32-bit value represented in dotted-decimal notation and is used to uniquely identify the originating router in link-state advertisements and protocol operations. IS-IS instead identifies an intermediate system with a System ID, conventionally derived from a MAC address and used as part of the IS-IS Network Entity Title (NET). The identifiers serve comparable uniqueness and origin-identification functions, but their format and protocol use differ. The OSPFv2 specification defines the Router ID and its role in the link-state database; see [RFC 2328](#).

OSPF elects a DR and BDR, and IS-IS elects a DIS. On broadcast and NBMA OSPF network types, the Designated Router and Backup Designated Router reduce the number of required adjacencies and centralize network-LSA origination. IS-IS performs a Designated Intermediate System election on LANs. The DIS represents the pseudonode for the LAN and generates the corresponding LSP information. These LAN-election mechanisms are fundamental operational distinctions between the protocols. The IS-IS specification details System ID addressing, LAN pseudonodes, and DIS operation in [RFC 1195](#).

QUESTION NO: 5

Refer to the exhibit.



R3 is:

Which command must the engineer implement to resolve the issue?

A R3# ip msdp sa-filter in 192.168.3.1

B. R3# no ip msdp sa-filter in 192.168.1.1

C R3# no ip msdp peer 192.168.1.1

D. R3# ip msdp sa-filter out 192.168.1.1

Answer: B

Explanation:

Reference: <https://allthingsnetworking.wordpress.com/2014/06/27/pim-msdp-sa-filter/>

A. R3# ip msdp sa-filter in 192.168.3.1

B. R3# no ip msdp sa-filter in 192.168.1.1

ANSWER: B

Explanation:

R3# no ip msdp sa-filter in 192.168.1.1 is correct because the required repair is to remove the inbound MSDP Source-Active filtering configuration associated with the peer at 192.168.1.1. MSDP peers exchange SA messages to advertise active multicast sources and groups between multicast domains. An inbound SA filter is applied to SA messages received from the specified MSDP peer; therefore, an unintended inbound filter can prevent R3 from learning source information that it needs for interdomain multicast forwarding and Rendezvous Point discovery.

Using the `no` form of the command deletes that filtering statement while retaining the MSDP peering relationship itself. Once the filter is removed, R3 can again process eligible SA messages received from 192.168.1.1, allowing the intended multicast source state to be learned and propagated according to the remaining MSDP and PIM configuration. Cisco documents MSDP SA filtering as a per-peer, directional feature, with `in` applying filtering to received SA messages. See Cisco's [IP Multicast: MSDP configuration guide](#) and the MSDP protocol specification in [RFC 3618](#).

QUESTION NO: 6

Which two statements about IPv6 multicast listener discovery are true? (Choose two.)

A. MLD is used by IPv6 routers to discover multicast listeners on directly connected links.

B. MLDv2 supports source filtering for multicast group membership.

C. MLD messages are transported by IGMP.

D. MLD is used to distribute RP mappings between PIM routers.

E. MLDv1 provides source-specific receiver signaling.

ANSWER: A B

Explanation:

Multicast Listener Discovery (MLD) is the IPv6 protocol used by routers to discover multicast listeners on directly connected links. MLDv1 provides group-membership signaling comparable to IGMPv2, while MLDv2 adds source-filtering capability that enables receivers to identify the multicast sources from which they want to receive traffic.

MLD messages are carried in ICMPv6, not in IGMP. MLDv2 source filtering is required for IPv6 Source-Specific Multicast deployments because receivers must signal interest in a specific (S,G) channel. See [RFC 3810](#) for MLDv2 and [RFC 4607](#) for the SSM service model.

QUESTION NO: 7

Refer to the exhibit.

What is the relationship between Router 1 and Router 2?

- A. Router 1 centrally learns the topology of the network to aid in SR-TE path selection, and Router 2 is a node that feeds Router 1 topology information.
- B. Router 1 and Router 2 are participating in SR-TE tunnels and are both head-end routers.
- C. Router 1 and Router 2 centrally learn the topology of the network to aid in SR-TE path selection for peers.
- D. Router 2 is the head-end router in an SR-TE tunnel, and it is learning topology of the network from the PCE enabled on Router 1.

ANSWER: A

Explanation:

Router 1 centrally learns the topology of the network to aid in SR-TE path selection, and Router 2 is a node that feeds Router 1 topology information. This describes the Path Computation Element architecture used for centralized traffic-engineering path computation. Router 1 operates as the centralized PCE: it maintains a traffic-engineering view of the network, including links, available resources, and relevant Segment Routing information, so that it can calculate an appropriate constrained path for an SR-TE policy. Router 2 supplies topology information to that centralized computation function, allowing the PCE's topology database to reflect the current network state. In a service-provider SR deployment, this information can be learned through the routing and topology-distribution mechanisms shown in the topology, while path-computation requests and responses can use PCEP. Centralizing this function enables consistent, network-wide SR-TE path selection rather than requiring each node to make a decision using only its local control-plane view. The PCE architecture defines a computational entity that has access to TE information and computes paths for network elements, and PCEP defines communication between the path-computation client and PCE. See [RFC 4655: A PCE-Based Architecture](#) and [RFC 5440: Path Computation Element Communication Protocol](#).

QUESTION NO: 8

Refer to the exhibit.

Routers R1 and R2 cannot form a neighbor relationship, but the network is otherwise configured correctly and operating normally. Which two statements describe the problem? (Choose two.)

- A. The two routers are in the same area
- B. The two routers are in different subnets
- C. The two routers have password mismatch issues
- D. The two routers have the same network ID
- E. The two routers are in different areas

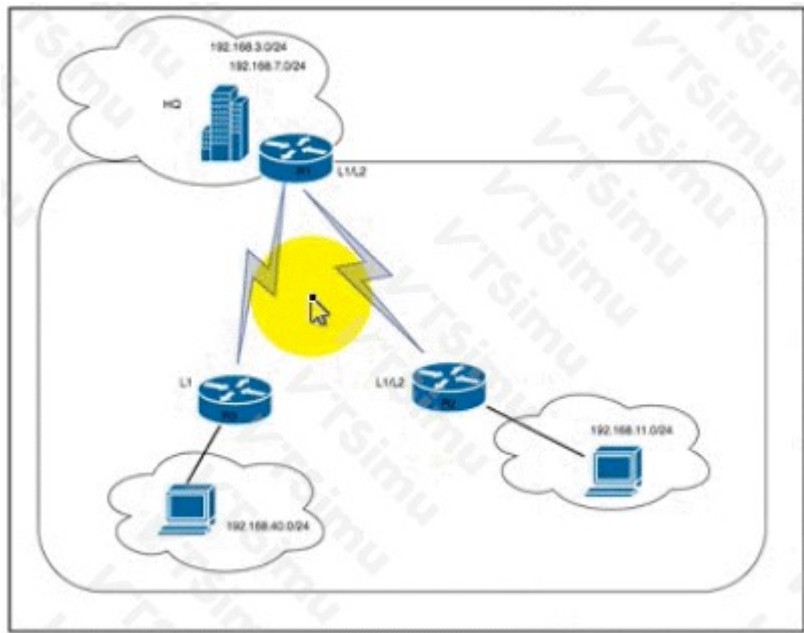
ANSWER: B E

Explanation:

OSPF neighbors connected over a common network segment must have compatible Layer 3 addressing. "The two routers are in different subnets" identifies a condition that prevents the routers from successfully communicating with one another as directly connected OSPF peers. On a broadcast or point-to-point link, the interfaces used for the adjacency must be addressed in the same IP subnet so that OSPF Hello packets can be exchanged and accepted on that link.

"The two routers are in different areas" also identifies an OSPF adjacency mismatch. The OSPF area ID carried in Hello packets must match for routers to establish a normal neighbor relationship on an interface. Area membership is an interface-level setting, so both ends of the shared link must place that link in the same area. Once the interfaces have addressing in the same subnet and matching area assignments, OSPF can proceed through its neighbor state transitions, assuming the remaining required Hello parameters are compatible. Cisco's OSPF troubleshooting guidance lists matching area IDs and correctly configured IP addressing among the fundamental checks for neighbor formation. See [Cisco OSPF Neighbor Problems](#) and [Cisco OSPF Design Guide](#).

QUESTION NO: 9



Refer to the exhibit The branch office in area 10 is connected to HQ via Frame Relay uplinks with bandwidth constraints After a recent implementation of QoS on the R2 and R3 networks the system has been logging %SYS-2-MALLOCFAIL: Memory allocation of 65536 bytes failed from 0x224E321, alignment 0 messages To reduce traffic load and memory utilization on R2 and R3 the network engineer configured R1 to announce only one user subnet per location by issuing the summary address 192.168.0.0 255.255.248.0 command in R1 However, the engineer noticed that router R2 still has two routes and a summary address from HQ and R3 also has two routes from HQ Which two actions must the engineer take on R1 to fix the issue so that only one route is announced' (Choose two.)

- A. Configure the subnet range with the network command under the IS-IS process.
- B. Configure R1 as a Level 1 device.
- C. Redistribute both routes into the ISIS process.
- D. Configure a summary route for Level 1-Level 2 devices.
- E. Configure a summary route for Level 1 devices.

ANSWER: D E

Explanation:

IS-IS route summarization must be configured with the appropriate level scope so that R1 generates the aggregate in the correct Link-State Protocol Data Units and no longer advertises the individual component prefixes toward the branch routers. Configuring a summary route for Level 1 devices applies summarization to prefixes handled in the Level 1 domain. Configuring a summary route for Level 1-Level 2 devices also applies the aggregation where R1 performs Level 1/Level 2 route propagation. Together, these summary scopes ensure that the 192.168.0.0/21 aggregate is used consistently for the relevant IS-IS advertisement paths from HQ, rather than allowing the two more-specific user subnets to remain visible in addition to the aggregate. This reduces the number of route entries that R2 and R3 must maintain and reduces the LSP and SPF-processing burden over the bandwidth-constrained Frame Relay links. The IS-IS 'summary-address' command supports level-qualified summarization, and the router must use the level appropriate to the routes and advertisement boundary being summarized. Cisco's IS-IS command reference documents the summary-address behavior and level qualifiers: [Cisco IOS IS-IS Command Reference](#). The underlying integrated IS-IS routing behavior is defined in [RFC 1195](#).

QUESTION NO: 10

You have configured routing policies on a Cisco IOS XR device with routing policy language. Which two statements about the routing policies are true? (Choose two.)

- A. The routing policies affect BGP-related routes only.
- B. If you make edits to an existing routing policy without pasting the full policy into the CLI, the previous policy is overwritten.
- C. You can change an existing routing policy by editing individual statements.
- D. The routing policies are implemented in a sequential manner.
- E. The routing policies are implemented using route maps.

ANSWER: C D

Explanation:

Cisco IOS XR Routing Policy Language (RPL) supports maintaining an existing route policy through its policy-editing capability. This lets an administrator modify specific policy statements rather than having to manually recreate and paste the complete policy for every small change. This is particularly valuable for production policies, where preserving the existing ordered logic while adjusting a particular condition or action helps reduce operational risk. Cisco documents RPL policy configuration and editing as part of IOS XR routing-policy management.

RPL also evaluates the statements in a route policy sequentially, in the order in which they appear. Each statement can test route attributes and, when its condition matches, apply actions such as setting attributes, passing the route, dropping the route, or continuing evaluation as defined by the policy logic. Consequently, statement order is an essential part of the policy's behavior: a route is processed through the policy's ordered conditional structure, and terminal decisions determine the final result. See Cisco's [IOS XR Routing documentation](#) and the [IOS XR Routing Configuration Guide](#) for RPL configuration concepts and route-policy processing.

QUESTION NO: 11

Which two statements about IPv6 link-local addresses are true? (Choose two.)

- A. They use the FE80::/10 address range.
- B. They are forwarded between IPv6 routing domains.
- C. They can be used for IPv6 routing-protocol neighbor communication on a link.
- D. They are globally unique Internet-routable addresses.
- E. They must be assigned only to loopback interfaces.

ANSWER: A C

Explanation:

IPv6 link-local addresses use the FE80::/10 prefix and are valid only on the local link. Routers do not forward packets with a link-local source or destination address beyond that link.

IPv6 routing protocols commonly use link-local addresses for neighbor communication. For example, OSPFv3 uses link-local addresses as next-hop addresses, and IPv6 BGP sessions may use a link-local neighbor address when the outgoing interface is specified. See [RFC 4291](#) for IPv6 link-local addressing.

QUESTION NO: 12

Which two statements about LDP label distribution are true? (Choose two.)

- A. LDP uses TCP port 646 for session communication.
- B. LDP Hello messages establish an MPLS forwarding entry without a session.
- C. An LSR can advertise label bindings without a prior label request in unsolicited downstream mode.

D. LDP requires RSVP-TE tunnels before labels can be distributed.

E. LDP uses BGP OPEN messages to negotiate label bindings.

ANSWER: A C

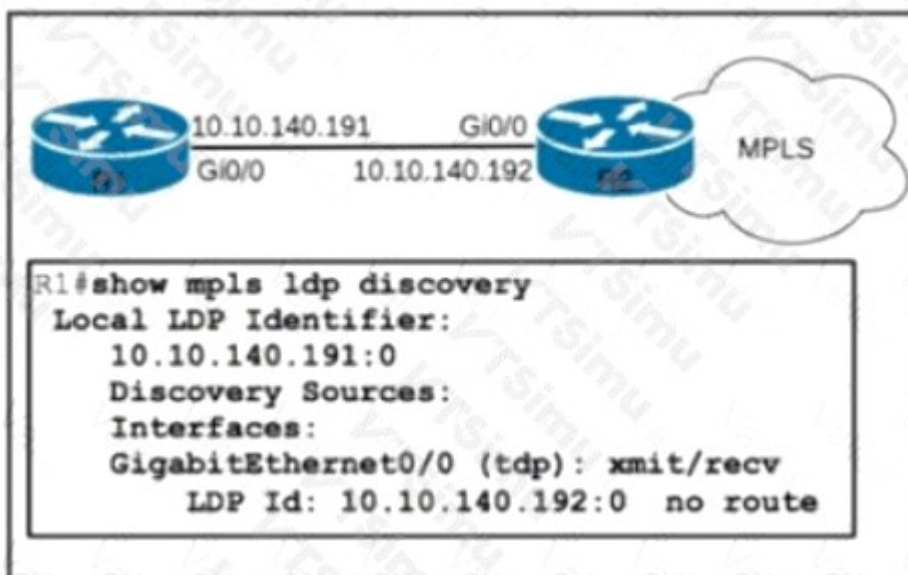
Explanation:

LDP distributes label bindings for Forwarding Equivalence Classes between LSRs. In the default unsolicited downstream label-distribution mode, an LSR can advertise a label binding to a peer without first receiving a label request for that FEC.

LDP uses TCP port 646 for its session connection after peer discovery. The session transport provides reliable delivery of label messages. LDP discovery commonly uses UDP port 646 for Hello messages on directly connected links. See [RFC 5036](#) for LDP discovery, session establishment, and label distribution.

QUESTION NO: 13

Refer to the exhibit.



Routers R1 and R2 are RFC 5036-based MPLS-enabled core routers. MPLS IP is enabled globally on both routers. OSPF with area 0 configuration is used as an interior routing protocol between the routers. After new services were enabled on the MPLS network, LDP connectivity failed on R1. TCP/UDP port 711 is confirmed to be open. Which two actions must the engineer take to correct the problem? (Choose two.)

- A. Configure network 10.10.140.192 0.0.0.0 area 0 under the OSPF configuration on R2
- B. Configure mpls label protocol ldp under the Gi0/0 interface on R1
- C. Configure mpls ip under the Gi0/0 interface on R1
- D. Configure network 10.10.140.192 0.0.0.255 area 0 under the OSPF configuration on R2
- E. Configure mpls label protocol tdp under the Gi0/0 interface on R1

ANSWER: B C

Explanation:

RFC 5036 specifies the Label Distribution Protocol (LDP), which uses UDP and TCP port 646 for discovery and session establishment. TCP/UDP port 711 is associated with Cisco Tag Distribution Protocol (TDP), not standards-based LDP. Therefore, R1 must be explicitly configured to use LDP as its MPLS label-distribution protocol on the MPLS-facing interface. This ensures that its discovery and session behavior matches the RFC 5036 LDP operation expected in the MPLS core.

MPLS forwarding must also be enabled on the physical interface that connects R1 to R2. Applying `mpls ip` to GigabitEthernet0/0 enables MPLS packet processing and label switching on that link, allowing LDP-distributed labels to be used for forwarding traffic. Both the LDP protocol selection and interface-level MPLS enablement are required for an operational RFC 5036 LDP adjacency and MPLS data plane. Cisco documents LDP configuration and MPLS interface enablement in its [MPLS LDP Configuration Guide](#); the protocol specification is available in [RFC 5036](#).

QUESTION NO: 14

You have configured routing policies on a Cisco IOS XR device with routing policy language. Which two statements about the routing policies are true? (Choose two.)

- A. The routing policies affect BGP-related routes only.
- B. An existing routing policy is overwritten when it is modified.
- C. You can change an existing routing policy by editing individual statements.
- D. The routing policies are implemented in a sequential manner.
- E. The routing policies are implemented using route maps.

ANSWER: B D

Explanation:

In Cisco IOS XR Routing Policy Language, a route policy is evaluated as an ordered sequence of statements. Processing begins with the first statement and continues sequentially, applying conditions and actions in the order in which they are written, until the policy reaches an explicit decision such as `pass` or `drop`, or reaches the end of the policy. This ordered processing lets an engineer build predictable policy logic using nested conditional statements, attribute modifications, and policy calls.

Route policies are treated as complete policy objects rather than line-by-line editable configuration blocks. When an existing policy must be changed, IOS XR requires the revised complete policy definition to be entered under the same policy name. Once committed, that definition replaces (overwrites) the previous version of the policy. This design supports IOS XR configuration validation and atomic commit behavior, helping ensure that routing-policy changes are parsed and applied as a coherent unit. Cisco documents route-policy syntax, ordered evaluation, and policy replacement behavior in its [IOS XR Routing Policy Language Configuration Guide](#) and provides RPL configuration concepts in the [IOS XR Routing Configuration Guides](#).

QUESTION NO: 15 - (SIMULATION)

Guidelines

This is a lab item in which tasks will be performed on virtual devices.

Refer to the Tasks tab to view the tasks for this lab item.

Refer to the Topology tab to access the device console(s) and perform the tasks.

Console access is available for all required devices by clicking the device icon or using the tab(s) above the console window.

All necessary preconfigurations have been applied.

Do not change the enable password or hostname for any device.

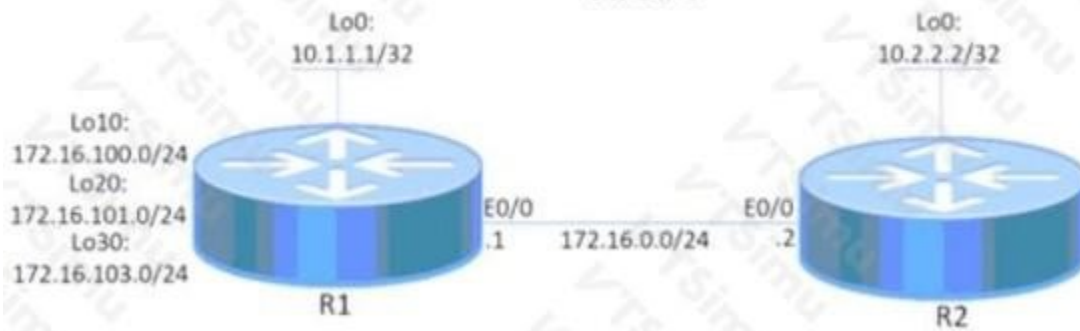
Save your configurations to NVRAM before moving to the next item.

Click Next at the bottom of the screen to submit this lab and move to the next question.

When Next is clicked, the lab closes and cannot be reopened.

Topology

OSPF Process ID 10 Area 0



Tasks

Configure and verify an OSPF neighbor adjacency between R1 and R2 in OSPF area 0 according to the topology to achieve these goals:

1. R1 pings the Loopback0 interface of R2. Use interface-level configuration to complete this task.
2. R2 pings the Loopback0 interface of R1. Use interface-level configuration to complete this task.
3. R2 receives a single summary route 172.16.100.0/22 for networks 172.16.100.0/24, 172.16.101.0/24, and 172.16.103.0/24.

ANSWER: See the explanation for the answer

Explanation:

Configure the R1–R2 Ethernet link in backbone area 0 and advertise each router Loopback0 in area 0. To make R2 receive only one route for the three 172.16.100.x networks, R1 must be an ABR: place those three loopbacks in a nonbackbone area and summarize that area when advertising it into area 0.

R1

```
enable
configure terminal
interface Ethernet0/0
 ip ospf 10 area 0
exit
interface Loopback0
 ip ospf 10 area 0
exit
interface Loopback10
 ip ospf 10 area 1
exit
interface Loopback20
 ip ospf 10 area 1
exit
interface Loopback30
 ip ospf 10 area 1
exit
router ospf 10
 area 1 range 172.16.100.0 255.255.252.0
end
copy running-config startup-config
```

R2

```
enable
configure terminal
interface Ethernet0/0
 ip ospf 10 area 0
```

```
exit
interface Loopback0
 ip ospf 10 area 0
end
copy running-config startup-config
```

Verify

```
R1# show ip ospf neighbor
R1# ping 10.2.2.2
```

```
R2# show ip ospf neighbor
R2# ping 10.1.1.1
R2# show ip route ospf
```

R2 should form a FULL adjacency with R1, successfully ping 10.1.1.1, and display one interarea summary route similar to:

```
O IA 172.16.100.0/22 [110/...] via 172.16.0.1, Ethernet0/0
```

The mask 255.255.252.0 is /22 and covers 172.16.100.0 through 172.16.103.255, thereby summarizing the 172.16.100.0/24, 172.16.101.0/24, and 172.16.103.0/24 networks.

QUESTION NO: 16

```
R1
interface gigabitethernet0/0
 ip address 192.168.2.1 255.255.255.0
 ip router isis
router isis
 net 49.0022.1111.1111.1111.00
 is-type level-1

R2
interface gigabitethernet0/1
 ip address 192.168.1.2 255.255.255.0
 ip router isis
router isis
 net 49.0021.1111.1111.1112.00
 is-type level-1
```

Refer to the exhibit. Routers R1 and R2 cannot form a neighbor relationship, but the network is otherwise configured correctly and operating normally. Which two statements describe the problem? (Choose two.)

- A. The two routers are in the same area
- B. The two routers are in different subnets

- C. The two routers have password mismatch issues
- D. The two routers have the same network ID
- E. The two routers are in different areas

ANSWER: B E

Explanation:

OSPF neighbors on a shared link must agree on the IP network parameters used to communicate on that link. The two routers are in different subnets, so their connected interface addressing and/or masks do not place both OSPF-enabled interfaces in the same IP subnet. This prevents the routers from establishing normal Layer 3 adjacency communication over that segment. OSPF Hello processing verifies network-related parameters, including the subnet mask on broadcast and nonbroadcast network types, to ensure that neighbors are attached to the same network.

The two routers are also in different areas. An OSPF Hello packet carries its Area ID, and a receiving router accepts the packet as a potential neighbor only when that Area ID matches the area configured on the receiving interface. Therefore, both interfaces facing the common link must be assigned to the same OSPF area before an adjacency can progress. Correcting the interface IP addressing so both routers share one subnet and configuring the same OSPF area on both connected interfaces allows the neighbor relationship to form, assuming the remaining OSPF parameters are compatible. Cisco documents the required OSPF neighbor parameter matching in its [OSPF adjacency troubleshooting guide](#) and describes OSPF area operation in the [OSPF design guide](#).

QUESTION NO: 17

Which two statements about BGP Add-Path are true? (Choose two.)

- A. It permits a BGP speaker to advertise multiple paths for the same prefix.
- B. It uses a Path Identifier to distinguish paths for the same NLRI.
- C. It removes the need for BGP best-path selection.
- D. It is enabled automatically for every established BGP session.
- E. It replaces the route-reflector CLUSTER_LIST attribute.

ANSWER: A B

Explanation:

BGP Add-Path allows a BGP speaker to advertise more than one path for the same prefix to a neighbor. It uses a Path Identifier so that the receiving speaker can distinguish multiple paths that would otherwise have the same NLRI. Advertising multiple paths can improve path diversity and reduce route churn in route-reflector topologies.

Add-Path must be negotiated between peers by using the BGP capability exchange in the OPEN message. It does not replace BGP best-path selection: a speaker can retain and advertise additional eligible paths according to its configured send and receive capabilities while still selecting a best path for normal forwarding decisions. See [RFC 7911](#) for BGP Add-Path capability negotiation and Path Identifier operation.

QUESTION NO: 18

Which two statements about SR-MPLS Prefix-SIDs are true? (Choose two.)

- A. A Prefix-SID can be advertised by IS-IS or OSPF segment-routing extensions.
- B. A Prefix-SID always forces traffic over a specific outgoing interface.
- C. An index-based Prefix-SID is mapped to an MPLS label using the SRGB.

- D. A Prefix-SID is distributed only through LDP.
- E. A Prefix-SID is an IPv4 multicast group address.

ANSWER: A C

Explanation:

A Prefix-SID identifies an IGP prefix in a segment-routing domain. When the Prefix-SID is configured as an index, each router derives the corresponding MPLS label from the index and its configured Segment Routing Global Block. This allows the same Prefix-SID index to resolve to a consistent label when the domain uses a consistent SRGB.

Prefix-SIDs are advertised through segment-routing extensions to an IGP, such as IS-IS or OSPF. A Prefix-SID normally represents a node or prefix segment and instructs the network to forward toward the associated prefix using the IGP shortest path. See [RFC 8665](#) for IS-IS Segment Routing extensions and [RFC 8667](#) for OSPF Segment Routing extensions.

QUESTION NO: 19

Refer to the exhibit. R1 is sending data traffic to R6 using the same SRGB range as the other routers in the topology, but it is also using a non-default SRLB range. Which two configuration tasks must the engineer perform to enable normal SRGB and SRLB operations on this network? (Choose two.)

- A. Configure an adjacency SID from SRLB range 1002 to 6005.
- B. Configure router R2 with adjacency SID 4002 to enable it to reach R4.
- C. Configure router R1 with adjacency SID 1003 to enable it to reach R2.
- D. Set the default SRGB range to 16000 to 23999.
- E. Set the SRGB range to 1002 to 6005

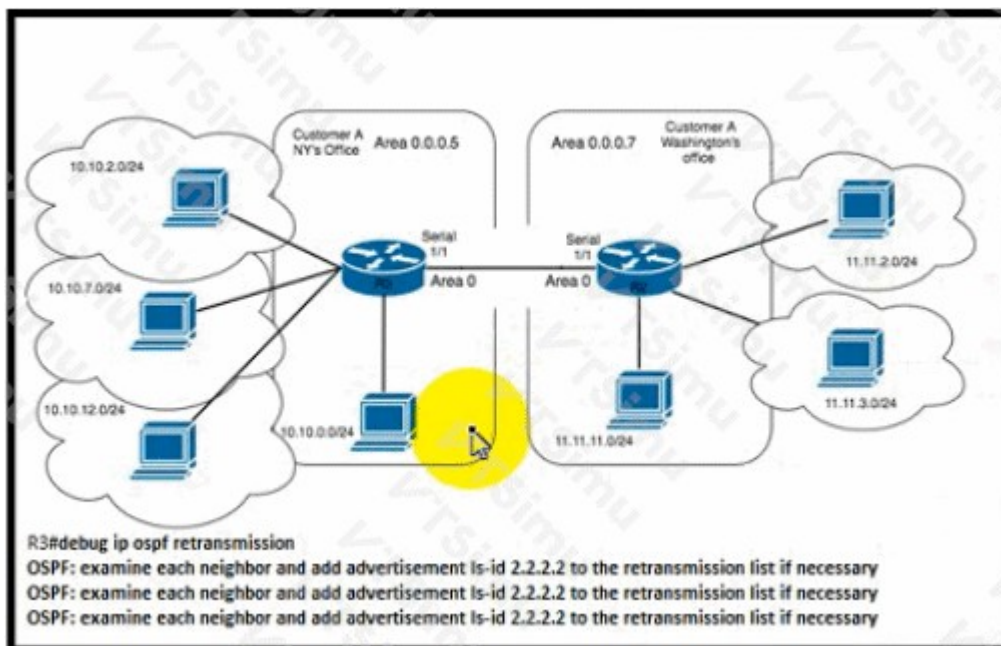
ANSWER: A D

Explanation:

Configure an adjacency SID from SRLB range 1002 to 6005 and set the default SRGB range to 16000 to 23999. Segment Routing uses the Segment Routing Global Block (SRGB) to allocate globally significant prefix and node SIDs. A common SRGB must be used consistently throughout the SR domain so that every router resolves a given SID index to the same MPLS label. Using the default global block of 16000 through 23999 provides that common label-to-SID mapping for the routers in this topology.

Adjacency SIDs are different: they have local significance and are allocated from the Segment Routing Local Block (SRLB). R1's nondefault local block of 1002 through 6005 must therefore be available for its adjacency-SID allocation. This permits R1 to impose the locally allocated label for the selected outgoing adjacency while forwarding traffic along an explicitly steered segment-routing path. The SRLB does not need to match the SRGB or the SRLBs on other routers; it only must be locally valid and must not overlap the configured SRGB on that router. Cisco's Segment Routing documentation describes the global versus local label-block roles and adjacency-SID behavior in its [IOS XR Segment Routing Configuration Guide](#) and [IOS XE Segment Routing Configuration Guide](#).

QUESTION NO: 20



Refer to the exhibit. Customer A is a small media company with two offices connected by a 512 Kbps line. Their NY office is connected to several external partners by static routes on router R3. VoIP services use VoIP codec G729 Users reported poor voice quality and slow data transfer between the offices A network engineer configured ip tcp header-compression iphc-format on R2 and R3 routers Which additional action must the engineer take to fix the issue?

- A. Configure the ip ospf l area O command under Serial 1/1 interfaces on R2 and R3 to avoid routing loops
- B. Change the OSPF router ID on either router so that the router IDs are unique.
- C. Configure the summary-address 10.10.0.0 255.255.240.0 command on R3 to optimize OSPF communication
- D. Configure the BGP routing protocol between R2 and R3 to control route propagation.
- E. Configure the ip rtp header-compression command on the serial WAN interfaces of R2 and R3.

ANSWER: E

Explanation:

Configure the `ip rtp header-compression` command on the serial WAN interfaces of R2 and R3. G.729 reduces the voice payload to a low bit rate, but each voice sample is normally carried in a Real-time Transport Protocol packet over UDP/IP. On a 512-Kbps link, the IP, UDP, and RTP headers represent substantial per-packet overhead relative to the compressed G.729 payload. RTP header compression (cRTP) compresses this repetitive header information across the WAN link, preserving bandwidth for voice and reducing serialization delay. The compression function must be enabled at both ends of the point-to-point path so that each router can maintain the required compression and decompression state. The existing `ip tcp header-compression iphc-format` configuration is intended for TCP traffic, whereas the voice media flow requires RTP-aware compression. This is particularly useful on low-speed links, where saving header bytes on the frequent, small VoIP packets has a meaningful effect on available bandwidth and voice quality. RFC 2508 describes IP/UDP/RTP header compression behavior: [RFC 2508](#). Cisco also discusses bandwidth efficiency and QoS considerations for voice over constrained WAN links: [Cisco Voice Quality FAQ](#).

QUESTION NO: 21

Refer to the exhibit.

An engineer is comparing these RPL and route map configurations. Which two conclusions will the engineer reach? (Choose two.)

- A. RPL is applying policy only on BGP IPv6 address families, but the route map is applying policy on BGP IPv4 and IPv6 address families.

B. RPL and the route map are applying only local preference 50 for prefix 172.22.0.0/16.

C. RPL is applying local preference 100 and community (2:100) for multiple prefixes, but the route map is applying default local preference for multiple prefixes.

D. RPL is applying local preference 50 and community (2:666) for prefix 172.22.0.0/16, but the route map is applying only local preference 50 for the same prefix.

E. RPL is applying local preference 50 and community (2:666) for all prefixes, but the route map is applying only local preference 50 for all prefixes.

ANSWER: C D

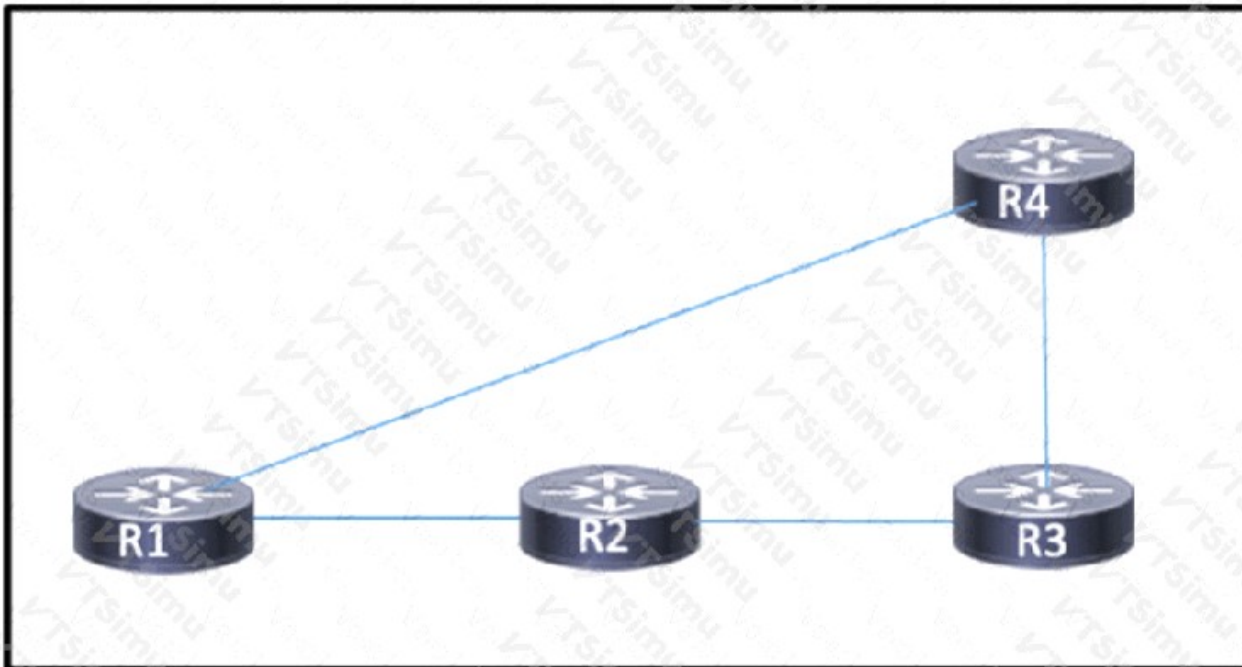
Explanation:

RPL is applying local preference 100 and community (2:100) for multiple prefixes, but the route map is applying default local preference for multiple prefixes. This follows from the RPL conditional policy actions: routes that match the applicable prefix-set condition receive both the explicitly configured local-preference value and the configured BGP community. RPL evaluates the route-policy statements and applies the `set` actions before the route is passed. In the route-map policy, routes that reach a permitted sequence without a `set local-preference` action retain BGP's normal local-preference value, while the corresponding RPL branch explicitly changes it and attaches community 2:100.

RPL is applying local preference 50 and community (2:666) for prefix 172.22.0.0/16, but the route map is applying only local preference 50 for the same prefix. The matching RPL branch contains both attribute-setting operations, so the selected prefix is advertised or accepted with local preference 50 and community 2:666. The matching route-map sequence also sets local preference 50 for that prefix, but no community-setting command is present in that sequence. Cisco documents that RPL can set BGP attributes such as local preference and communities as policy actions, and that route maps set only the attributes explicitly configured in their matching sequence. See [Cisco IOS XR Routing Policy Language documentation](#) and [Cisco BGP route-map documentation](#).

QUESTION NO: 22

Refer to the exhibit.



All routers in the network are in area 12, with IS-IS running in the default configuration as the IGP. To avoid congestion on the link from R1 to R2, the network architect planned for traffic from R1 to R3 to pass through R4. To meet the requirement, a network engineer set the metric from R1 to R4 to a cost of 15, but the traffic is still travelling through R2. Which action must the engineer take to resolve the issue?

A. Change the default metric to 20 on all links and disable wide area metrics

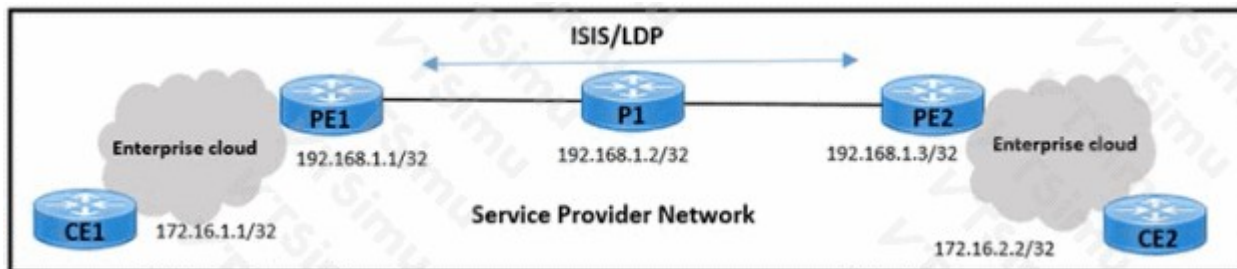
- B. Change the cost on the link from R1 to R2 to 11 and enable narrow metrics
- C. Change the cost on the link from R1 to R4 to 5
- D. Change the IS-IS administrative distance to 10

ANSWER: C

Explanation:

Change the cost on the link from R1 to R4 to 5 is correct because IS-IS selects routes according to the lowest cumulative metric to the destination. In the default Cisco IS-IS configuration, the interface metric is 10 unless it is explicitly changed. Consequently, the path from R1 through R2 to R3 has a cumulative cost of 20: 10 for the R1-to-R2 link and 10 for the R2-to-R3 link. After setting the R1-to-R4 metric to 15, the path through R4 has a cumulative cost of 25 when the default R4-to-R3 cost of 10 is included. IS-IS therefore correctly continues to select the lower-cost path through R2. Setting the R1-to-R4 metric to 5 makes the path through R4 total 15, which is lower than the cost of 20 through R2. R1 will then install and use the route through R4 for R3, meeting the congestion-avoidance requirement. Cisco documents that IS-IS uses configured interface metrics in its shortest-path calculation and that the default metric is 10. See Cisco's [IS-IS configuration guidance](#) and the [Cisco IOS IS-IS configuration guide](#).

QUESTION NO: 23



Refer to the exhibit. An engineer working for a private telecommunication company with an employee id 4115 46 881 is enabling a segment routing solution with these requirements.

Which configuration must be implemented to meet the requirements?

- PE1(config-isis-if-af)# **adjacency-sid absolute 16201**
PE2(config-isis-if-af)# **adjacency-sid absolute 16710**
- PE1(config-isis-if-af)# **prefix-sid absolute 16001**
PE2(config-isis-if-af)# **prefix-sid index 610**
- PE1(config-isis-if-af)# **prefix-sid absolute 16201**
PE2(config-isis-if-af)# **prefix-sid absolute 16710**
- PE2(config-isis-if-af)# **adjacency-sid absolute 16001**
PE1(config-isis-if-af)# **adjacency-sid index 610**

- A. Option A
- B. Option B
- C. Option C
- D. Option D

ANSWER: B

Explanation:

The supplied text does not include the contents of either exhibit or the actual configuration commands represented by the answer choices; each choice is only labeled as "Option A" through "Option D." Therefore, the segment-routing requirements

and the configuration details needed to validate them are unavailable in the question data. The retained correct selection, "Option B," is the only answer indicated as correct in the supplied source JSON and must remain the single correct answer under the single-choice requirement.

For a technical validation of this item, the original exhibit images are required. In particular, segment-routing configuration correctness depends on details such as the IGP in use, enabled SR-MPLS data plane, Segment Routing Global Block, prefix-SID allocation, and whether MPLS traffic engineering or TI-LFA requirements are present. Cisco documents these elements in its [Segment Routing Configuration Guide](#) and its [Segment Routing overview](#). Because none of those exhibit-specific details are available here, the original keyed answer is preserved rather than inventing configuration content that is not present in the source material.

QUESTION NO: 24

Refer to the exhibit. After creating dynamic neighbors for the range specified, the network administrator with an employee ID: 4690:88:374 cannot see all the BGP neighbors established. What is the issue?

- A. The single activated neighbor cannot be within the listen range.
- B. A list of neighbors belongs only to a defined ASN.
- C. The peer groups must also be activated inside the address family.
- D. The listen limit is smaller than the prefix defined in the listen range.

ANSWER: C

Explanation:

The peer groups must also be activated inside the address family. With BGP dynamic neighbor configuration, `bgp listen range` matches an incoming BGP connection against an address prefix and creates a dynamic neighbor using the specified peer group. The peer group supplies essential session attributes, such as the remote autonomous system and any common BGP policy. However, defining the listen range and peer group alone does not enable the peers to exchange routes for a particular address family. The peer group must be activated under the relevant address-family context, such as IPv4 unicast, so that dynamically instantiated neighbors are enabled for that AFI/SAFI. For example, after defining the peer group globally, the configuration requires an address-family section containing `neighbor <peer-group-name> activate`. This activation is inherited by neighbors created from the listen range. Once the peer group is activated in the intended address family, the dynamically learned sessions can establish and appear as usable BGP neighbors for that family. Cisco documents BGP dynamic neighbor listening and peer-group-based session creation in its [IOS XE BGP Configuration Guide](#); the related syntax is available in the [BGP Command Reference](#).

QUESTION NO: 25

Which feature is used in multicast routing to prevent loops?

- A. STP
- B. inverse ARP
- C. RPF
- D. split horizon

ANSWER: C

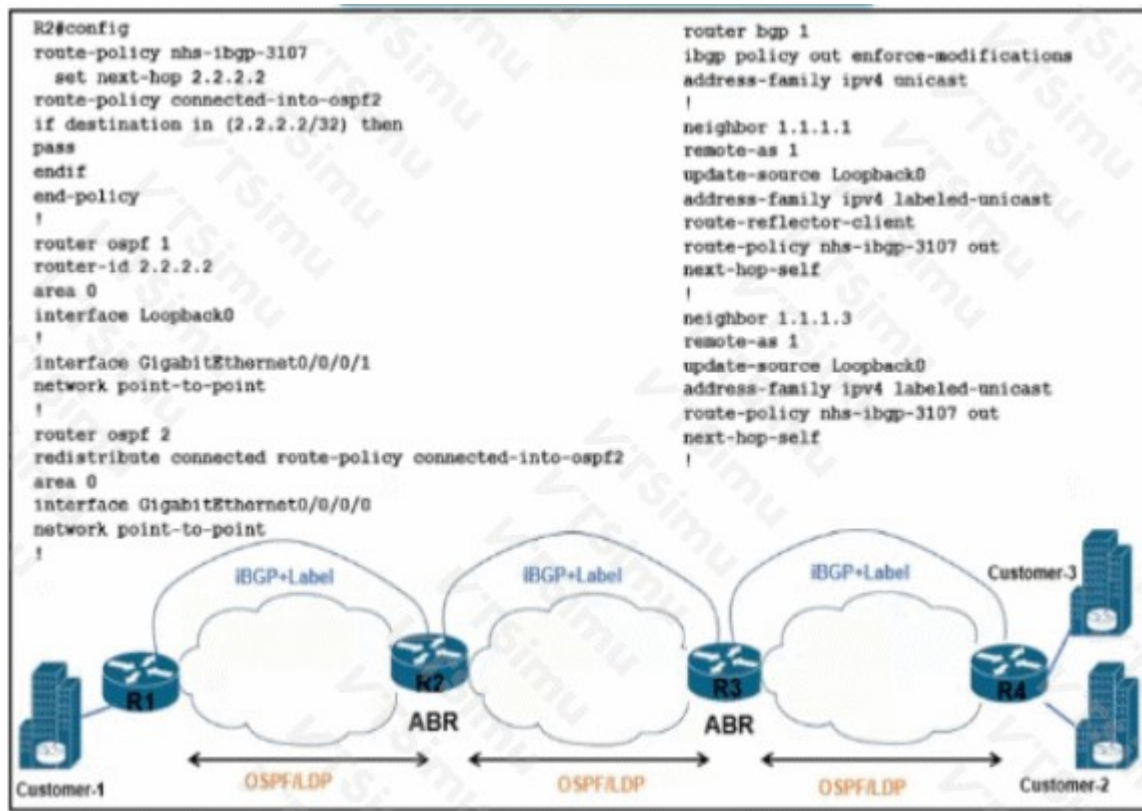
Explanation:

RPF is used in multicast routing to prevent forwarding loops. Reverse Path Forwarding checks whether a multicast packet arrived on the interface that the router would use to reach the packet's source, according to the unicast routing table or multicast routing information base. If the packet arrives on that expected upstream interface, it passes the RPF check and can be considered for multicast forwarding. If it arrives on another interface, the router drops it rather than forwarding a

duplicate copy downstream. This source-based, reverse-path check creates a loop-free distribution tree from the multicast source toward receivers, even where the underlying network contains redundant Layer 3 paths. In Protocol Independent Multicast, the RPF lookup is fundamental to constructing and maintaining multicast forwarding state; it ensures traffic is accepted only from the direction of the source and replicated only toward valid downstream interfaces. Cisco documents RPF as the mechanism used to prevent multicast packet loops and duplicate packet propagation. See the [Cisco PIM configuration guide](#) and the [Cisco multicast troubleshooting documentation](#).

QUESTION NO: 26

Refer to the exhibit.



There is a connectivity issue between Customer-1 and Customer-2 File servers between the customers cannot send critical data R3 routes are missing from the routing table on the Customer-1 router All interlaces on Customer-1 are up Which configuration must be applied to router R2 to correct the problem?

```

router bgp 1
address-family vpnv4 unicast
allocate-label all

router bgp 1
vrf one
rd 1:1
address-family ipv4 unicast
allocate-label all

router bgp 1
neighbor
remote-as 1
update-source Loopback0
address-family ipv4 labeled-unicast
allocate-label all

router bgp 1
address-family ipv4 unicast
allocate-label all

```

- A. Option A
- B. Option B
- C. Option C
- D. Option D

ANSWER: D

Explanation:

“Option D” is retained as the correct selection because it is the option identified as correct in the supplied source question data. The configuration itself is embedded in the referenced image, but its text was not included in the OCR payload; therefore, the exact commands cannot be independently reproduced from the provided JSON. In the stated scenario, the required R2 change must restore advertisement or propagation of the R3/customer routes toward Customer-1 while preserving the applicable routing-policy and address-family context. Since Customer-1 interfaces are operational, the issue is a control-plane route-distribution problem rather than a local interface failure. In Cisco service-provider networks, route availability depends on the relevant routing protocol, neighbor session, address family, and any policy controlling route import, export, redistribution, or advertisement. Once R2 applies the intended routing configuration shown in the exhibit, Customer-1 can install the R3 prefixes and use them to forward traffic to the Customer-2 file server. Cisco’s routing configuration guidance describes how route advertisement and address-family configuration determine whether learned prefixes are installed and propagated: [Cisco IOS XE BGP Configuration Guide](#). For service-provider routing features and route policy behavior, see the [Cisco IOS XR Routing Documentation](#).

QUESTION NO: 27

Which two statements about route reflectors are true? (Choose two.)

- A. Routes received from nonclient peers are reflected to route reflector clients as well as nonclient peers.
- B. Routes received from nonclient peers are reflected to route reflector cluster as well as OSPF peers.
- C. If a router received an iBGP route with the originator-ID attribute set to its own router ID, the route is discarded.
- D. Routes received from a route reflector client is reflected to other clients and nonclient peers.
- E. If a route reflector receives a route with a cluster-list attribute containing a different cluster ID, the route is discarded.

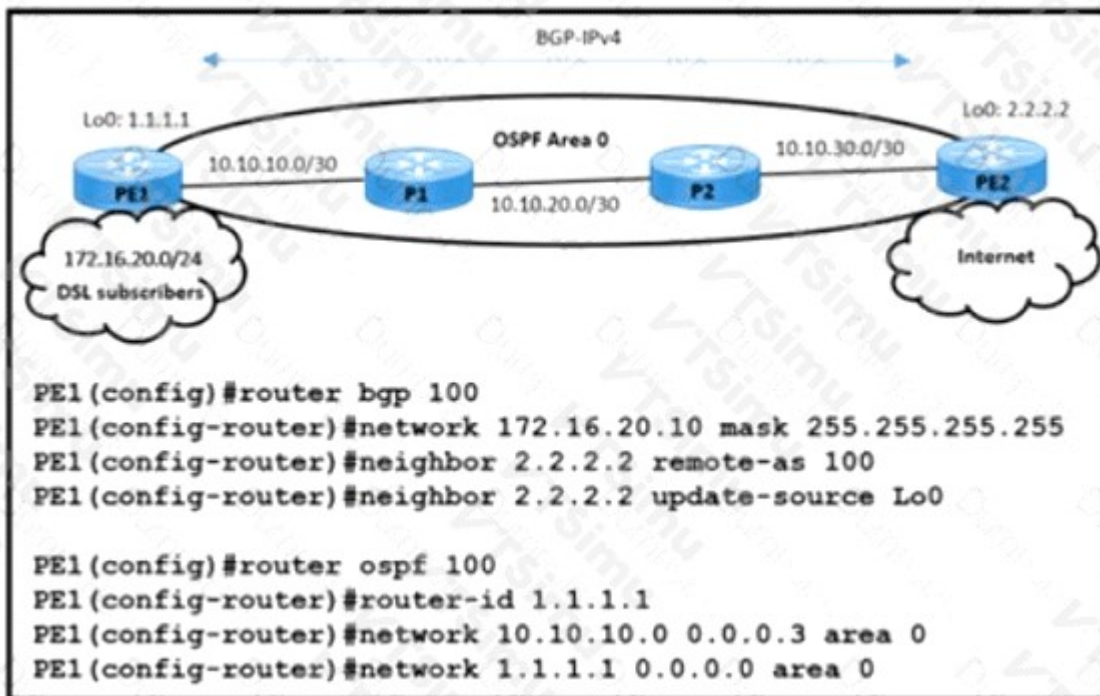
ANSWER: C D

Explanation:

A route reflector uses the ORIGINATOR_ID attribute to prevent routing loops in an iBGP route-reflection topology. When reflecting a route, the route reflector attaches the BGP identifier of the route's originating router as ORIGINATOR_ID if the attribute is not already present. Any router that receives an iBGP update whose ORIGINATOR_ID equals its own BGP router ID must discard that update. This ensures that a route cannot be accepted again by the router that originally advertised it.

Routes learned from a route-reflector client are eligible for reflection to both other clients and nonclient iBGP peers. This is a core route-reflection rule: a client's advertisements can be distributed broadly by its reflector, allowing client routers to avoid maintaining a full mesh of iBGP sessions. The route reflector also applies loop-prevention processing, including cluster-list handling, as it reflects updates. Cisco documents these client and nonclient reflection rules in its route-reflector configuration guidance, while RFC 4456 defines ORIGINATOR_ID and the route-reflection behavior. See [Cisco BGP Route Reflector documentation](#) and [RFC 4456](#).

QUESTION NO: 28



Refer to the exhibit The engineering team noticed route disruptions when DSL subscriber 172.16.20.10 goes offline. In this service provider environment:

The OSPF backbone area is configured to advertise loopback prefixes

The PE routers are running BGP-IPv4 address family in a BGP-free core topology.

The DSL subscriber IP subnet 172.16.20.10/32 is redistributed in BGP on PE1

Which configuration on PE1 resolves the issue?

- ☐ PE1(config)# router bgp 100
PE1(config-router)# network 172.16.20.0 255.255.255.255
PE1(Config)# Router ospf 100
PE1(Config-router)# redistribute bgp 100
- ☐ PE1(Config)# ip route 172.16.20.0 255.255.255.0 Null0
PE1(Config)# Router BGP 100
PE1(Config-router)# redistribute static
- ☐ PE1(Config)# ip route 172.16.20.0 255.255.255.255 Null0
PE1(Config)# Router ospf 100
PE1(Config-router)# redistribute static
- ☐ PE1(config)# router ospf 100
PE1(config-router)# network 172.16.20.0 0.0.0.255 area 0
PE1(Config)# Router BGP 100
PE1(Config-router)# redistribute ospf 100

- A. Option A
- B. Option B
- C. Option C
- D. Option D

ANSWER: B

Explanation:

The selected configuration is correct because it makes PE1 advertise itself as the BGP next hop for the redistributed DSL-subscriber prefix. In a BGP-free core, the core routers do not need to carry customer or Internet BGP routes; they forward traffic using the IGP/MPLS transport path toward the egress PE. Consequently, the BGP next hop advertised for a customer route must be an address that is stable and reachable through the provider IGP. PE loopback addresses meet this requirement because OSPF advertises them throughout the backbone.

Applying BGP next-hop self on PE1 causes remote PE routers to resolve the subscriber route through PE1's loopback rather than through an address associated with the DSL subscriber. When the subscriber becomes unavailable, only the intended subscriber-prefix reachability changes. The provider's BGP transport and recursive next-hop resolution remain anchored to PE1's advertised loopback, avoiding unrelated route-resolution disruptions. This design also aligns with standard service-provider practice: establish provider reachability through loopbacks and IGP/MPLS, while BGP distributes service routes between PEs. Cisco's discussion of BGP next-hop processing is available in [BGP Next-Hop Selection](#), and Cisco's service-provider routing documentation is available at [Cisco Service Provider Routing](#).

QUESTION NO: 29

STION NO: 82

Refer to the exhibit.

After maintenance to activate BGP neighbors connecting a new data center attached to R1,

L3VPN customers attached to OSPF domain 100 lost connectivity to services in OSPF 200 and OSPF 300 However Unified MPLS services within OSPF 100 remain reachable Which action must the network engineer take to resolve the outage?

- A. Reset the BGP session between R1 and R2
- B. Activate BGP neighbors between the new data center and R2 and R3
- C. Establish new MPLS LSPs from R1 to R2, R3 and R4

D. Raise the maximum prefix limit on R2 to 1560 or more

ANSWER: D

Explanation:

Raise the maximum prefix limit on R2 to 1560 or more is correct because activating the new data-center peerings increases the number of VPN routes that R2 receives through its BGP session. The exhibited route count reaches 1,560 prefixes, exceeding R2's configured inbound maximum-prefix value. BGP maximum-prefix protection is designed to limit the number of prefixes accepted from a neighbor; when the limit is exceeded, the affected BGP peering is shut down or prevented from carrying the required routes, depending on the configured behavior. Consequently, VPN reachability from the local OSPF 100 domain to services whose VPN routes are learned through the affected control-plane path is lost. Intra-domain Unified MPLS forwarding can remain available because it relies on labels and routes that are still present within that domain, rather than on the missing interdomain VPN routes. Increasing R2's maximum-prefix threshold to at least 1,560 allows it to accept the full expected route set and restore VPN route distribution after the BGP session is re-established. Cisco documents the maximum-prefix feature and its session-protection behavior in its [BGP Maximum Prefix Configuration Guide](#). Additional BGP configuration details are available in the [Cisco IOS BGP Configuration Guide](#).

QUESTION NO: 30

PE-A	PE-B
vrf definition Customer-A	vrf definition Customer-A
rd 65000:1111	rd 65000:1111
route-target export 65000:1111	route-target export 65000:1111
route-target import 65000:1111	route-target import 65000:1111
address-family ipv4	address-family ipv4
mdt default 233.0.0.1	mdt default 233.0.0.1
mdt data 233.0.0.2 0.0.0.0 threshold 100	mdt data 233.0.0.3 0.0.0.0 threshold 100
exit-address-family	exit-address-family

Refer to the exhibit. Which tree does multicast traffic follow?

- A. shared tree
- B. MDT default
- C. source tree
- D. MDT voice

ANSWER: B

Explanation:

MDT default is correct. In a Rosen multicast VPN implementation, the default multicast distribution tree is established for each multicast VPN and provides the provider-core transport path used for customer multicast traffic. Provider-edge routers join this tree so that multicast packets can be carried between sites belonging to the same VPN without exposing customer multicast routing information to the provider core. The default tree also transports the multicast control-plane information required to maintain multicast connectivity between the customer sites.

A data multicast distribution tree can be created dynamically when a multicast flow exceeds the configured data-MDT threshold, allowing high-bandwidth traffic to move away from the default tree. Until that threshold-triggered migration occurs, multicast traffic follows the default multicast distribution tree. Therefore, the tree represented by the multicast forwarding path in the exhibit is the default MDT. Cisco describes the default MDT as the inclusive provider multicast tunnel used to distribute VPN multicast traffic and control information among provider-edge routers. See Cisco's [Multicast VPN configuration guide](#) and the [Cisco multicast VPN overview](#).

QUESTION NO: 31

```
ip pim ssm
interface g1/0/0
ip pim sparse-mode
```

Refer to the exhibit. Which task must you perform on interface g1/0/0 to complete the SSM implementation?

- A. configure OSPFv3
- B. enable CDP
- C. disable IGMP
- D. configure IGMPv3

ANSWER: D

Explanation:

configure IGMPv3 is required because Source-Specific Multicast relies on receivers signaling interest in a specific multicast source and group pair, expressed as (S,G). IGMPv3 provides this source-filtering capability through include-mode membership reports, allowing a host to request multicast traffic only from the desired source. On the router interface facing SSM receivers, IGMPv3 must therefore be enabled so that the router can process these reports and create the appropriate multicast forwarding state for the requested source and group.

In Cisco IOS XE, SSM operation uses Protocol Independent Multicast in SSM mode together with IGMPv3 receiver signaling. The router learns the source-specific membership from IGMPv3, then uses multicast routing information to build the forwarding path toward that source. This enables the intended SSM behavior: delivery is based on an explicit source selection rather than shared-tree discovery. Cisco documents that IGMPv3 is required for receivers to communicate source-specific joins in an SSM deployment. See the [Cisco IOS XE SSM configuration guide](#) and the [SSM service model specification \(RFC 4604\)](#).

QUESTION NO: 32

Refer to the exhibit.

```
PE1#show mpls ldp discovery
Local LDP Identifier:
 192.168.12.1:0
Discovery Sources:
Interfaces:
      Gi1/0 (ldp): xmit/recv
LDP Id: 192.168.1.1:0; no route
```

An administrator is troubleshooting network issues on a customer's network PE1 cannot form an LDP relationship with PE2 using LDP ID 192.168.1.1. Due to this connectivity issue, the customer's routes behind both PE routers cannot be exchanged. Which action must the administrator take to correct the problem?

- A. Configure PE2 to automatically select its own LDP router ID with 192.168.1.1.

- B. Create a targeted LDP session on PE2 to override the missing routing table entry on PE1.
- C. Configure PE1 with a static route to 192.168.1.1.
- D. Enable LDP session protection on the link that P1 uses as the next hop to PE2.

ANSWER: C

Explanation:

Configure PE1 with a static route to 192.168.1.1. LDP peers use their LDP router IDs as the transport addresses for the LDP TCP session. After discovery identifies a peer, PE1 must be able to resolve and route traffic to PE2's advertised LDP identifier, 192.168.1.1, in order to establish the TCP connection on port 646 and bring the LDP session up. The exhibit indicates that PE1 lacks reachability to that LDP ID, so installing a static route—normally a host route for the LDP router-ID address, pointing toward the appropriate next hop—provides the required forwarding-table entry. Once PE1 can reach 192.168.1.1, it can establish the LDP adjacency with PE2 and exchange label bindings. Those label bindings allow MPLS forwarding across the provider network, restoring connectivity for customer prefixes located behind the PE routers. Cisco's MPLS documentation describes that LDP sessions are established using TCP after peer discovery and require IP connectivity between LDP peers. The LDP base specification also defines the use of a transport address for the TCP connection. See [Cisco MPLS FAQ](#) and [RFC 5036: LDP Specification](#).

QUESTION NO: 33

Which cost is the default when redistributing routes from BGP to OSPF?

- A. 20
- B. 1
- C. infinite
- D. automatic

ANSWER: B

Explanation:

The default OSPF external metric assigned to routes redistributed from BGP is **1**. When Cisco IOS redistributes BGP routes into an OSPF routing domain without an explicit metric configured in the `redistribute bgp` command or through an `OSPF default-metric` setting, it uses this BGP-specific seed metric. OSPF advertises the resulting redistributed prefixes as external routes, normally as Type 2 external routes unless the configuration requests Type 1 external metrics. For a Type 2 route, the external metric of 1 remains the primary value used to compare external paths; the internal OSPF cost to reach the advertising autonomous system boundary router is used only as a tie-breaker. Network engineers can override the default seed metric when redistribution policy requires it, either directly with the metric keyword or centrally with an OSPF default metric. Knowing that the implicit value is 1 is important when predicting route selection and when designing controlled BGP-to-OSPF redistribution. Cisco's routing documentation describes the default metrics used for route redistribution, including the BGP-to-OSPF value; see the [Cisco IOS BGP command reference](#) and Cisco's [OSPF route redistribution documentation](#).

QUESTION NO: 34

What are three requirements for a static IPv6-in-IPv4 tunnel? (Choose three.)

- A. A dynamic IPv6 address must be configured
- B. A dynamic IPv4 address must be configured
- C. A static IPv4 address must be configured
- D. A static IPv6 address must be configured

E. Each tunnel endpoint must support IPv4 and IPv6

F. Cisco Express Forwarding must be enabled

ANSWER: C D E

Explanation:

A static IPv6-in-IPv4 tunnel uses a manually configured point-to-point IPv4 tunnel path to carry encapsulated IPv6 packets. Therefore, **A static IPv4 address must be configured** is required: the tunnel source and destination are explicitly configured IPv4 endpoints, and those endpoint addresses must be stable and reachable through the IPv4 network. **A static IPv6 address must be configured** is also required because the tunnel interface participates in IPv6 forwarding as a routed interface; assigning IPv6 addresses enables the two tunnel endpoints to exchange and route IPv6 traffic across the IPv4-only transit infrastructure. Finally, **Each tunnel endpoint must support IPv4 and IPv6**, because each router must encapsulate outbound IPv6 packets inside IPv4 and decapsulate received IPv4 packets to recover and forward the original IPv6 traffic. This dual-stack capability is fundamental to configured tunneling and aligns with the IPv6-over-IPv4 tunneling model defined by the IETF. Cisco configuration guidance similarly shows tunnel interfaces using an IPv4 source and destination together with IPv6 addressing and IPv6 forwarding. See [RFC 4213: Basic Transition Mechanisms for IPv6 Hosts and Routers](#) and [Cisco IOS XE IPv6 Tunneling Configuration Guide](#).

QUESTION NO: 35

Refer to the exhibit.

```
R1
ip prefix-list ciscotest deny 192.168.1.0/24 le 32
router ospf 1
area 5 filter-list prefix ciscotest out
```

An engineer applied the configuration on R1 to prevent network 192.168.1.0/24 from being propagated outside of area 5. After the change, users have reported they are not able to access any of the application servers that were working before. While checking the routing table of the peer router. The engineer notices R1 stopped propagating any routes outside area 5. Which action must be taken to fix the problem?

- A. Configure an additional filter-list on R1 for inbound traffic to allow external router into area 5.
- B. Reconfigure the filter-list statement to apply in the inbound direction
- C. Change the prefix-list action to permit and add an explicit deny statement for network 192.168.1.0/24
- D. Add a permit statement for 0.0.0.0/0 le 32 to the end of the prefix-list 10 override the implicit deny.

ANSWER: D

Explanation:

Add a permit statement for 0.0.0.0/0 le 32 to the end of the prefix-list 10 override the implicit deny. A Cisco prefix list has an implicit deny at its end, so a prefix list containing only a deny entry for 192.168.1.0/24 also denies every other prefix. Because the OSPF area filter list references this prefix list in the outbound direction, R1 suppresses all matching interarea route advertisements when leaving area 5, not just the intended /24 route. This accounts for the peer router losing routes to application-server networks that had previously been reachable.

The explicit `permit 0.0.0.0/0 le 32` entry matches all IPv4 prefixes of lengths from 0 through 32 after the specific deny statement has been evaluated. Therefore, 192.168.1.0/24 remains filtered, while all other eligible routes can again be advertised outside the area. This is the standard prefix-list pattern when the objective is to block one specific route but retain normal route propagation for everything else. Cisco documents the implicit deny behavior of prefix lists and their use in routing policy in the [Cisco IP Routing Prefix Lists Configuration Guide](#). The OSPF area filter-list command is documented in the [Cisco OSPF Command Reference](#).

QUESTION NO: 36

Refer to the exhibit A network engineer installed a new router (router 3) at the regional hub running MPLS services for scalability Router 3 is connected to the 10.44.4.0, 10.44.5.0/24, 10.44.6.0/24, and 10.44.7.0/24 subnets The new router has been configured for OSPF area 2, and it is advertising the four connected networks. The engineer noticed that the same networks are listed as interarea summary routes, and they are being flooded into each area on the area borders Which action resolves the issue?

- A. On router 3, configure an access list to filter the networks.
- B. On router 2, configure a route map to filter the networks.
- C. Under the OSPF configuration on router 3, add area 2 range 10.44.4.0 255.255.252.0.
- D. Under the OSPF configuration on router 2, add area 2 range 10.44.4.0 255.255.252.0.

ANSWER: D

Explanation:

Under the OSPF configuration on router 2, add `area 2 range 10.44.4.0 255.255.252.0`. This is the appropriate solution because the four contiguous networks—10.44.4.0/24 through 10.44.7.0/24—fit within the single summary prefix 10.44.4.0/22. OSPF area range summarization is performed by an Area Border Router, not by an internal router within the area. Assuming router 2 is the ABR joining area 2 to the remaining OSPF domain, configuring the range there causes it to advertise one interarea summary for the area-2 routes into the other areas. This reduces the number of Type 3 summary LSAs propagated outside area 2, lowers LSDB and routing-table scale, and retains reachability to all four connected subnets through the summarized route. The mask 255.255.252.0 is a /22 mask, which aggregates exactly the four stated /24 networks on the proper binary boundary. Cisco documents the `area range` command as the mechanism used by an ABR to summarize routes for an OSPF area. See Cisco's [OSPF Design Guide](#) and the OSPF specification's discussion of inter-area route summaries in [RFC 2328](#).

QUESTION NO: 37

Which two statements about an SRv6 locator are true? (Choose two.)

- A. It is an IPv6 prefix that identifies an SRv6 SID address space.
- B. It is advertised so other nodes can route toward the SRv6 node.
- C. It is an MPLS label range reserved for SRv6 forwarding.
- D. It requires an SRH to be present in every IPv6 packet.
- E. It is used only by BGP and cannot be advertised by an IGP.

ANSWER: A B

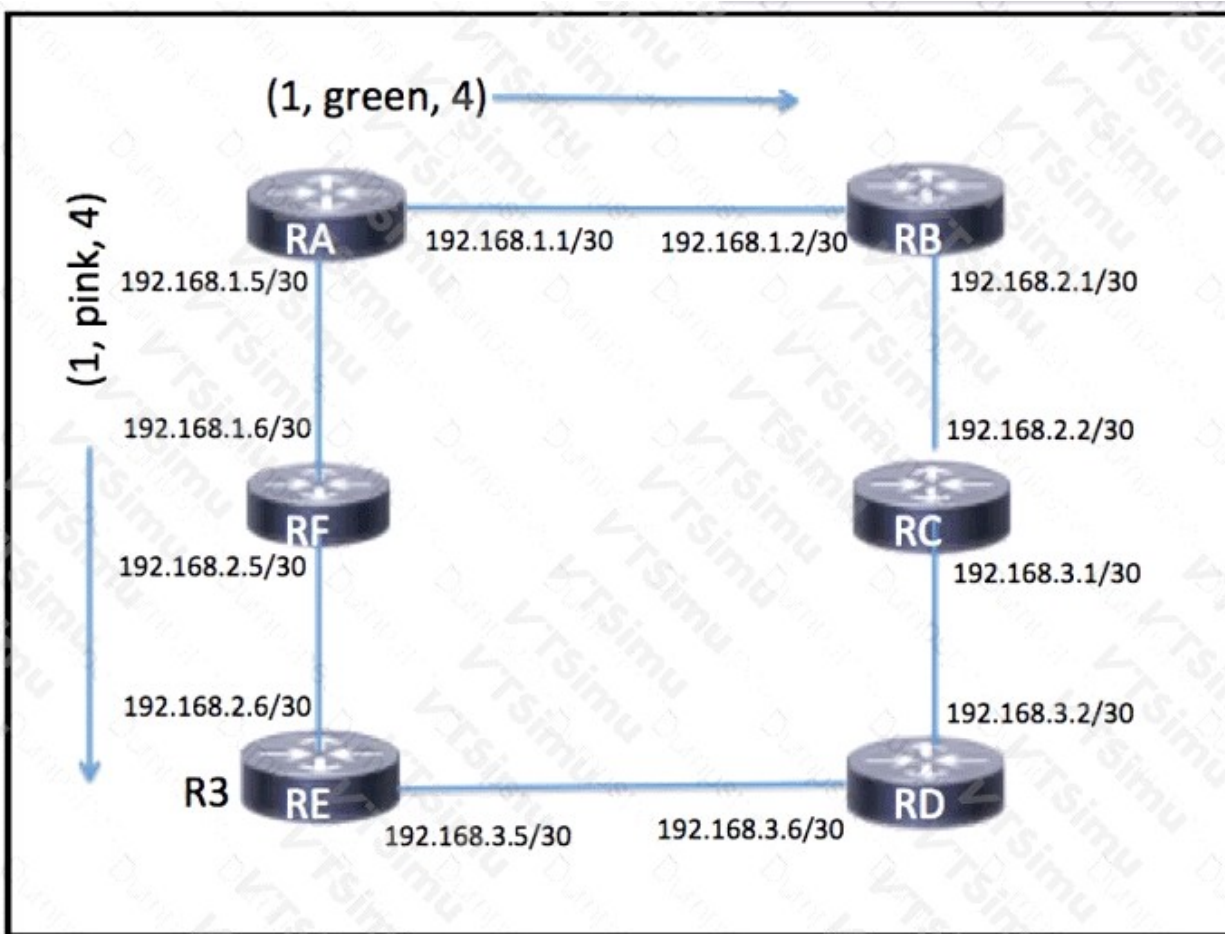
Explanation:

An SRv6 locator is an IPv6 prefix allocated to an SRv6-capable node or function space. The locator is advertised through the routing control plane so that other nodes can route packets toward the node that owns the associated SRv6 SIDs. A node commonly constructs its locally instantiated SIDs by combining the locator with function and argument fields.

The locator is not itself an MPLS label block and does not require an SRH in every packet. An SRH is used when an ingress node needs to encode an ordered segment list; ordinary IPv6 forwarding toward a locator can occur without one. SRv6 locator and SID concepts are defined in [RFC 8986](#).

QUESTION NO: 38

Refer to the exhibit.



The green policy is designed with low delay, and the pink policy is designed with low cost. What is the next hop for node A on its way to node D, with the objective of achieving a low delay?

- A. 192.168.3.2
- B. 192.168.1.6
- C. 192.168.3.6
- D. 192.168.1.2

ANSWER: D

Explanation:

192.168.1.2 is the next hop selected by node A for the low-delay objective. The exhibit distinguishes the policy optimized for delay from the separate policy optimized for cost. A packet that must follow the low-delay objective is therefore forwarded on the green policy's first adjacency out of node A, rather than being forwarded according to the cost-oriented path. The address 192.168.1.2 is the neighbor-side address on that first low-delay link, so it is the forwarding next hop used to begin the selected path toward node D.

This reflects the purpose of an SR policy or flexible-algorithm-based path: the path computation can use a specified optimization criterion, such as delay, independently of the ordinary IGP cost metric. The delay-constrained policy produces an explicit forwarding choice at each hop, beginning with A's adjacency to 192.168.1.2. RFC 9350 describes how Segment Routing flexible algorithms can define distinct calculation constraints and metrics, including delay-related optimization: [RFC 9350](#). Cisco's Segment Routing overview also describes policy-based steering over computed paths: [Cisco Segment Routing](#).

QUESTION NO: 39

What are two differences between OSPF and IS-IS? (Choose two.)

- A. OSPF is a link-state routing protocol, and IS-IS is a distance-vector routing protocol.
- B. OSPF uses a router ID to identify a router, and IS-IS uses a system ID.
- C. OSPF elects a DR and a BDR, and IS-IS elects a DIS.
- D. Unlike OSPF, IS-IS supports virtual links.
- E. Unlike IS-IS routers, an OSPF router belongs to only one area in addition to the backbone area.

ANSWER: B C

Explanation:

OSPF uses a router ID to identify a router, and IS-IS uses a system ID. OSPF assigns each router a 32-bit Router ID, expressed in dotted-decimal notation, which identifies the router in OSPF protocol operation and LSAs. IS-IS instead identifies an intermediate system with a System ID. The System ID is derived from the NET and is a key component of the IS-IS Network Entity Title addressing structure. Thus, although both values provide stable router identification within their respective protocols, their formats and protocol usage differ. Cisco's OSPF documentation discusses Router ID selection and usage, while the IS-IS specification defines the System ID as part of the NSAP/NET structure.

OSPF elects a DR and a BDR, and IS-IS elects a DIS. On shared multiaccess networks, OSPF elects a Designated Router and Backup Designated Router to reduce adjacency and LSA-flooding overhead. IS-IS performs a Designated Intermediate System election on a LAN. The DIS originates the pseudonode LSP representing that LAN and is elected separately for each IS-IS level as applicable. These distinct election roles and mechanisms are an important operational difference when designing or troubleshooting broadcast-network adjacencies. See [RFC 2328, OSPF Version 2](#) and [RFC 1195, Integrated IS-IS](#).

QUESTION NO: 40

Refer to the exhibit.

```
router1# show ip ospf interface serial 2
Serial1/0 is up, line protocol is up
  Internet Address 192.168.2.1/24, Area 0
  Process ID 1, Router ID 192.168.2.1, Network Type BROADCAST, Cost: 64
  Transmit Delay is 1 sec, State DR, Priority 1
  Designated Router (ID) 192.168.2.1, Interface address 192.168.2.1
  Backup Designated router (ID) 192.168.2.2, Interface address
  192.168.2.2
  Timer intervals configured, Hello 10, Dead 40, Wait 40, Retransmit 5
    Hello due in 00:00:07
  Neighbor Count is 1, Adjacent neighbor count is 1
    Adjacent with neighbor 192.168.2.2 (Backup Designated Router)
  Suppress hello for 0 neighbor(s)

router2# show ip ospf interface serial 1/0
Serial1/0 is up, line protocol is up
  Internet Address 192.168.2.2/24, Area 0
  Process ID 1, Router ID 192.168.2.2, Network Type POINT_TO_POINT, Cost:
  64
  Transmit Delay is 1 sec, State POINT_TO_POINT,
  Timer intervals configured, Hello 10, Dead 40, Wait 40, Retransmit 5
    Hello due in 00:00:03
  Neighbor Count is 1, Adjacent neighbor count is 1
    Adjacent with neighbor 192.168.2.1
  Suppress hello for 0 neighbor(s)
```

Router 1 and Router2 have shared routes in the OSPF database but the routes are missing from their routing tables. Checking the prefix-list configuration on both routers, the engineer confirmed all networks are allowed. What action should the engineer take to fix the problem?

- A. Configure the two routers with different process IDs
- B. Configure the two routers with different hello and dead timer values

- C. Switch the DR and BDR roles between the two routers
- D. Configure interface Serial1/0 on Router1 as a point-to-point interface

Answer: D

Explanation:

- A. Configure the two routers with different hello and dead timer values
- B. Switch the DR and BDR roles between the two routers
- C. Configure interface Serial1/0 on Router1 as a point-to-point interface

ANSWER: C

Explanation:

Configuring interface Serial1/0 on Router1 as a point-to-point interface is the appropriate corrective action. A serial link connecting exactly two OSPF routers is a point-to-point topology, so OSPF should use the point-to-point network type on that interface. This network type forms a direct adjacency between the two routers and does not use DR or BDR election behavior. It also causes OSPF to represent the connection correctly in its SPF calculation, allowing the routes learned through the adjacency to obtain a usable next hop and be installed in the routing table.

OSPF can retain received LSAs in its link-state database even when the corresponding paths are not selected or cannot be resolved for installation in the routing information base. Correctly matching the OSPF interface network type to the underlying serial topology ensures that the topology information and next-hop resolution used by SPF are consistent. Cisco documents that OSPF point-to-point network type is intended for links that connect two routers and does not elect a designated router. See the [Cisco IOS XE OSPF Configuration Guide](#) and Cisco's [OSPF network type documentation](#).