

Implementing and Operating Cisco Service Provider Network Core Technologies (350-501 SPCOR)

Cisco 350-501

Version Demo

Total Demo Questions: 68

Total Premium Questions: 681

Buy Premium PDF

<https://passqueen.com>

support@passqueen.com

Topic Break Down

Topic	No. of Questions
Topic 1, Architecture	35
Topic 2, Virtualization	22
Topic 3, Infrastructure	431
Topic 4, Network Assurance	50
Topic 5, Security	79
Topic 6, Automation	64
Total	681

QUESTION NO: 1

Refer to the exhibit.

A network engineer must meet these requirements to provide a connects, solution:

The customer connected to Area 2 needs to access the application in Area 1 on the 10.10.10.0/24 subnet

The Customer must not have access to the 20.10 30.0/24 subnet.

The service provider must make sure that the Area 2 routing database limits the number of IP addresses in the routing table

Which two configurations must be implemented to meet the requirements? (Choose two)

- A. Set a tag value of 200 to match the summary address 10.0.0/16 on R2.
- B. Set a tag value of 200 to match the summary address 10.0.0.0/16 on R3.
- C. Apply the route map for tag 200 and leak Level 2 routes into Level 1 Area 2 on R3
- D. Apply the route map for tag 200 and teak Level 2 routes into Level 1 Area 2 on R4.
- E. Set a tag value of 200 to match the summary address 10.0.0./16 on R1.

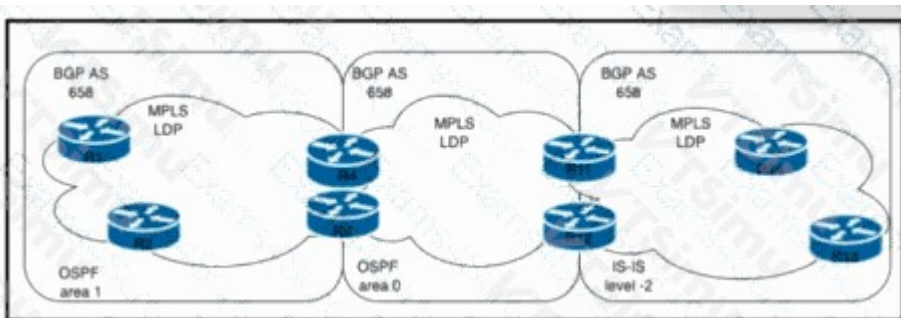
ANSWER: B C

Explanation:

“Set a tag value of 200 to match the summary address 10.0.0.0/16 on R3” and “Apply the route map for tag 200 and leak Level 2 routes into Level 1 Area 2 on R3” work together to selectively provide reachability while controlling the routes installed in Area 2. R3, as the relevant Level 1/Level 2 boundary router, creates the 10.0.0.0/16 summary and assigns it tag 200. The application prefix 10.10.10.0/24 falls within that aggregate, so advertising the summary gives Area 2 customers a route toward the application without introducing its individual subnet route into the Area 2 routing database.

The route map matches tag 200 during Level 2-to-Level 1 route leaking. Consequently, only the explicitly tagged 10.0.0.0/16 aggregate is leaked into Area 2. This provides a scalable route advertisement: Area 2 retains one summarized destination instead of multiple more-specific prefixes. The selective policy also ensures that only the intended summarized application address space is made available through the leak. IS-IS Level 1/Level 2 operation and inter-level route advertisement are defined in [RFC 1195](#). Cisco’s IS-IS configuration documentation describes using route-policy controls, including route maps, for routing-information handling: [Cisco IOS IS-IS Configuration Guide](#).

QUESTION NO: 2



Refer to the exhibit. ISP_A is a Tier-1 ISP that provides private Layer 2 and Layer 3 VPN services across the globe using an MPLS-enabled backbone network. Routers R1, R2, R13, and R14 are PE devices. ASRs R3 and R4 and ASBRs R11 and R12 are acting as route reflectors with preconfigured next-hop-self for all BGP neighbors. After the network experienced IGP RIB overflow with multiple routes, a network engineer is working on a more scalable network design to minimize the impact on local RIBs. The solution should be based on RFC 3107 and maintain existing QoS capabilities.

Which task must the engineer perform to achieve the goal?

- A. Implement Seamless Segment Routing on the ABR routers.

- B. Implement MPLS Traffic Engineering on PE routes only.
- C. Implement Hierarchical MPLS VPLS across all LDP-enabled routers.
- D. Implement Unified MPLS on all MPLS-enabled routers.

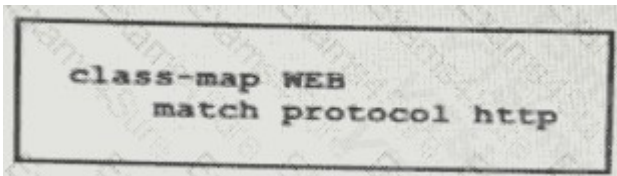
ANSWER: D

Explanation:

Implement Unified MPLS on all MPLS-enabled routers. Unified MPLS is designed to scale an MPLS transport network across multiple routing domains while avoiding the need to distribute large numbers of remote transport prefixes through the IGP. It uses BGP Labeled Unicast (BGP-LU), specified by RFC 3107, to advertise a prefix together with an MPLS label. This allows routers to establish labeled reachability for remote PE loopbacks and service endpoints through BGP, rather than requiring every such route to be present in each local IGP RIB. The existing route reflectors and ASBRs, with next-hop-self configured, are well suited to propagating this labeled reachability between network domains. Deploying Unified MPLS throughout the MPLS infrastructure provides an end-to-end label-switched transport design while retaining the existing Layer 2 and Layer 3 VPN services. MPLS QoS is also maintained because packet forwarding continues to use MPLS label stacks and their Traffic Class field, enabling preservation and application of QoS policy across the transport path. RFC 3107 defines the labeled NLRI encoding used for this BGP-based label distribution: [RFC 3107](#). Cisco's BGP labeled-unicast implementation guidance is available at [Cisco BGP Labeled Unicast Configuration Guide](#).

QUESTION NO: 3

Refer to the exhibit:



```
class-map WEB
  match protocol http
```

Which statement describes the effect of this configuration?

- A. It applies a service policy to all interfaces remarking HTTP traffic
- B. It creates an ACL named WEB that filters HTTP traffic.
- C. It matches HTTP traffic for use in a policy map
- D. It modifies the default policy map to allow all HTTP traffic through the router

ANSWER: C

Explanation:

The configuration matches HTTP traffic for use in a policy map. Cisco Modular QoS CLI uses a class map to define a traffic class by identifying packets that meet specified match criteria. When the configuration contains a class-map definition named `WEB` and a protocol-based HTTP match statement, packets identified as HTTP are placed into that class for subsequent QoS treatment. The class map itself is a classifier: it does not apply an action, attach a policy to an interface, or permit traffic. Instead, it provides a reusable traffic definition that can be referenced under a policy map with a `class` statement. The policy map can then associate actions such as marking, policing, shaping, bandwidth allocation, or queueing with the HTTP traffic class. Those actions take effect only when the completed policy map is applied in the appropriate direction on an interface or another supported policy attachment point. This separation of classification, policy definition, and policy attachment is a core feature of Cisco MQC. Cisco documents class maps as the mechanism for specifying match criteria used by policy maps; see the [Cisco MQC class-map configuration guide](#) and the [Cisco policy-map configuration guide](#).

QUESTION NO: 4

Refer to Exhibit.

```
username cisco privilege 15 password 0 cisco
!
ip http server
ip http authentication local
ip http secure-server
!
snmp-server community private RW
!
netconf-yang
netconf-yang cisco-ia snmp-community-string cisco
restconf
```

A network engineer is trying to retrieve SNMP MIBs with RESTCONF on the Cisco switch but fails. End-to-end routing is in place. Which configuration must the engineer implement on the switch to complete?

- A. netconf-yang cisco-ia snmp-community-string Public
- B. snmp-server community cisco RW
- C. snmp-server community public RO
- D. netconf-yang cisco-ia snmp-community-string Private

ANSWER: A

Explanation:

The required configuration is `netconf-yang cisco-ia snmp-community-string Public`. On Cisco IOS XE, access to SNMP MIB data through the model-driven management interfaces is provided through the Cisco IA (Cisco-IOS-XE-ia) mechanism. The `snmp-community-string` setting supplies the SNMP community string that the Cisco IA service uses when it obtains MIB information for NETCONF or RESTCONF requests. The configured value must correspond to the community string expected in the lab or exhibit, which is `Public` here. Once RESTCONF is enabled and the RESTCONF client can reach and authenticate to the device, this setting permits the IOS XE management framework to retrieve the requested SNMP MIB objects and return them in the RESTCONF response. This is separate from merely establishing IP reachability: the device must have the Cisco IA SNMP community configured so that the RESTCONF-to-SNMP MIB access can be performed. Cisco's IOS XE programmability documentation describes NETCONF/YANG and RESTCONF support at [Cisco Developer: NETCONF/YANG](#) and [Cisco Developer: RESTCONF](#).

QUESTION NO: 5

Why do Cisco MPLS TE tunnels require a link-state routing protocol?

- A. The link-state database provides segmentation by area, which improves the path-selection process.
- B. The link-state database provides a data repository from which the tunnel endpoints can dynamically select a source ID.
- C. Link-state routing protocols use SPF calculations that the tunnel endpoints leverage to implement the tunnel.
- D. The tunnel endpoints use the link-state database to evaluate the entire topology and determine the best path.

ANSWER: D

Explanation:

The tunnel endpoints use the link-state database to evaluate the entire topology and determine the best path. MPLS Traffic Engineering requires more than ordinary IP reachability information: the headend router must have a topology view and Traffic Engineering attributes for links, such as available reservable bandwidth, TE metrics, administrative groups (affinities),

and shared-risk link-group information where configured. Link-state IGPs, specifically OSPF or IS-IS with their TE extensions, distribute this information and maintain the Traffic Engineering Database (TED) used by the tunnel headend. The headend applies constrained shortest path first (CSPF) to that database to calculate an explicit route that satisfies the tunnel's configured constraints and signaling requirements. It then signals the selected path with RSVP-TE, reserving resources along the route as appropriate. This complete topology knowledge is essential because a valid TE path can depend on characteristics of every candidate link, not merely the next-hop route toward the tunnel destination. Cisco documents that TE uses extensions to link-state routing protocols to advertise TE information, and that CSPF uses the resulting TE topology to select tunnel paths. See Cisco's [MPLS Traffic Engineering configuration guide](#) and [Cisco documentation on MPLS TE tunnel path calculation](#).

QUESTION NO: 6

Refer to the exhibit. A company recently deployed a new network using OSPF in the core to share routes. The network administrator selected OSPF as the routing protocol because of its ability to maintain a route database. When the new network was started up, all routers booted normally, but the link between routers R1 and R2 failed to come up. The two routers are located in the same rack at the data center. Which task should an engineer perform to correct the problem?

- A. Synchronize the dead timers.
- B. Change one of the OSPF router IDs so that the router IDs are in different subnets
- C. Change the OSPF process ID on one of the devices so that the two IDs match
- D. Configure the MTUs on the interface to match.

ANSWER: D

Explanation:

Configure the MTUs on the interface to match. is correct because OSPF neighbors exchange Database Description packets while building adjacency and compare the interface MTU value carried in those packets. Each router must be able to accept packets up to the MTU advertised by its neighbor. If the connected interfaces use different MTU values, the routers can discover one another but fail during database synchronization, commonly leaving the adjacency in an ExStart or Exchange state rather than reaching Full. As a result, the routers cannot establish the OSPF neighbor relationship required to exchange and maintain the link-state database and routes.

The engineer should verify the Layer 3 MTU configured on both ends of the R1-to-R2 connection and make the values identical. This is especially important in data-center environments, where one interface may have been configured for jumbo frames while the peer remains at the default Ethernet MTU. Cisco documents OSPF MTU mismatch handling and notes that a mismatch can prevent adjacency establishment: [Cisco: OSPF Adjacency Problems](#). The relevant OSPF Database Description packet behavior is also defined in [RFC 2328, Section 10.6](#).

QUESTION NO: 7

What is the primary purpose of the flexible algorithm in a segment-routing environment?

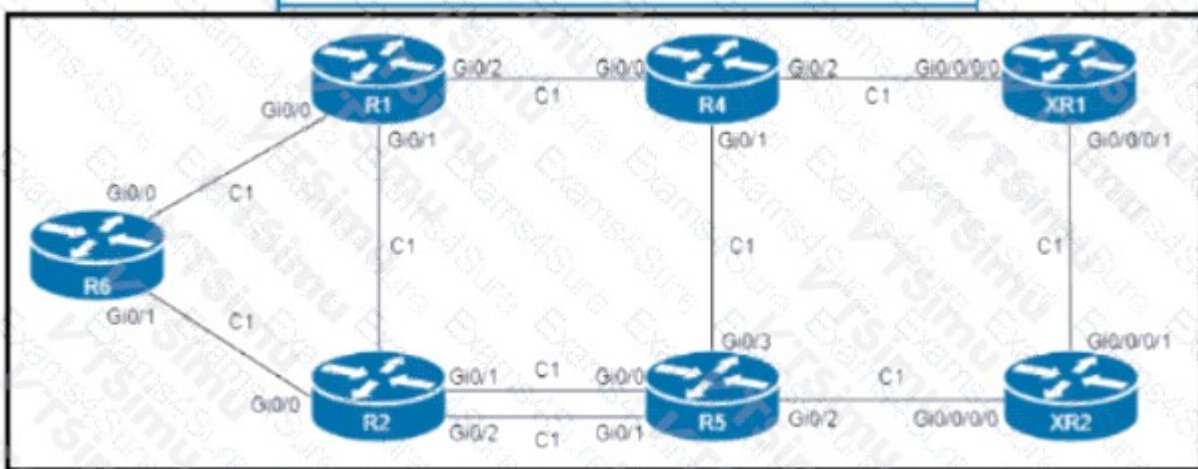
- A. to distribute designated-path information for certain traffic classes to maintain routing information in accordance with current network-performance indicators
- B. to integrate with legacy routing protocols to ensure backward compatibility and a smooth transition in mixed network environments
- C. to determine and assign labels for specific types of traffic, ensuring that routes maintain a valid backup path as network conditions change
- D. to support adaptive path selection, optimizing network performance based on real-time network metrics

ANSWER: A

Explanation:

The primary purpose of a flexible algorithm is to define and advertise a constraint-based IGP computation plane that produces separate Segment Routing paths for a particular traffic class or service requirement. Rather than relying only on the default shortest-path calculation, an operator can define an algorithm using a chosen metric and constraints, such as administrative-group (affinity) inclusion or exclusion, shared-risk link group exclusion, and metric type. Nodes advertise the flexible-algorithm definition and the corresponding algorithm-specific prefix SIDs through the IGP, allowing all participating routers to calculate paths consistently for that algorithm. This provides deterministic, policy-aligned forwarding paths that can be selected by applications, services, or traffic-engineering policies when their requirements differ from the regular IGP topology. Flexible Algorithm is therefore an SR control-plane capability for creating customized logical topologies and associated paths, while preserving distributed IGP path computation. Cisco documents its use for defining customized algorithm-based paths in Segment Routing deployments, and the protocol behavior is standardized in RFC 9350. See [Cisco IOS XR Segment Routing Flexible Algorithm documentation](#) and [RFC 9350](#).

QUESTION NO: 8



Refer to the exhibit. An engineer configured R6 as the headend LSR of an RSVP-TE LSP to router XR2, with the dynamic path signaled as R6-R2-R5-XR2, and set the OSPF cost of all links to 1. MPLS autotunnel backup is enabled on all routers to protect the LSP. Which two NNHOP backup tunnels should the engineer use to complete the implementation? (Choose two.)

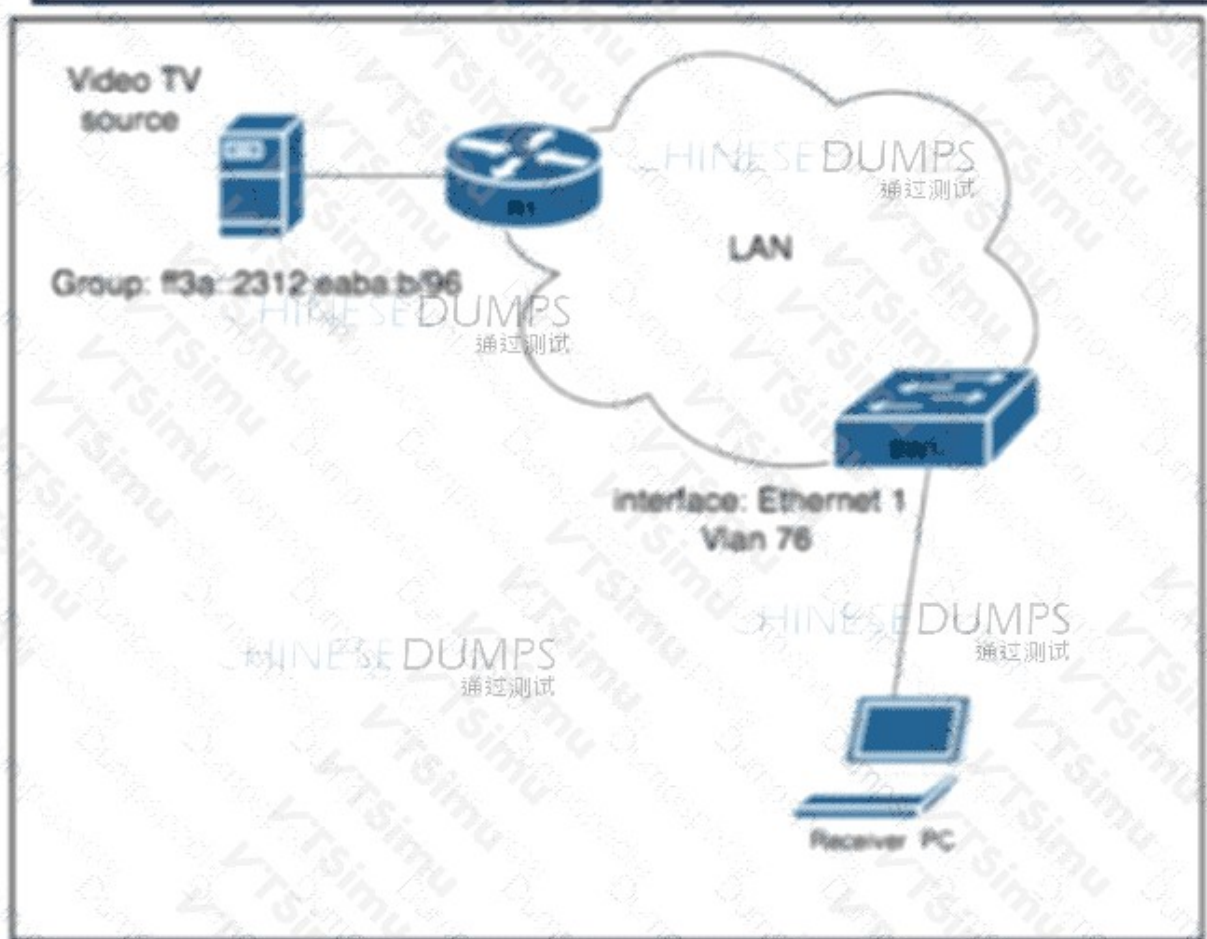
- A. The R6 backup tunnel path R6-R1-R4-R5.
- B. The R2 backup tunnel path R2-R5 across the alternate link.
- C. The R2 backup tunnel path R2-R1-R4-XR1-XR2.
- D. The R6 backup tunnel path R6-R2-R5
- E. The R6 backup tunnel path R6-R1-R2.

ANSWER: A C

Explanation:

Next-next-hop (NNHOP) protection uses RSVP-TE Fast Reroute facility backup tunnels to provide node protection. The point of local repair (PLR) creates a detour that bypasses its immediate downstream router and merges at a node farther along the protected LSP. For the primary LSP R6-R2-R5-XR2, R6 is the PLR for protection of R2. The backup tunnel path R6-R1-R4-R5 avoids R2 and merges at R5, the next-next hop on the original path. It therefore maintains a protected path toward XR2 if R2 or the R6-R2 adjacency fails.

R2 must likewise provide protection for its downstream next hop, R5. The backup tunnel path R2-R1-R4-XR1-XR2 bypasses R5 and merges at XR2, which is downstream of R5 on the working LSP. This is the appropriate node-protecting NNHOP detour from R2, because traffic can resume the original LSP forwarding context at the merge point without traversing the failed protected node. Autotunnel backup can automatically establish these TE backup tunnels when the appropriate FRR protection requirements are configured. Cisco describes RSVP-TE FRR backup tunnels and node protection in its [MPLS TE Fast Reroute documentation](#); the underlying facility-backup and merge-point behavior is also defined in [RFC 4090](#).



Refer to the exhibit. A network engineer working for a telecommunication company with an employee ID: 4602:62:646 is configuring security controls for the IPv6 multicast group, which is used for video TV. The solution from the engineer should reduce network usage and minimize the leave latency for the user that is connected to VLAN 76. Which two configurations meet this goal? (Choose two.)

A)

Apply the following commands globally on SW1:

ipv6 mld vlan 76 fast-leave vlan 76

ipv6 mld security join vlan 76

B)

Configure an ACL to limit the IPv6 multicast group with the entry **permit ipv6 any ff3a::2312:eaba:b/96**.

C)

Configure an ACL to limit the IPv6 multicast group with the entries **ipv6 access-list security_access_list** and **permit ipv6 ff3a::2312:eaba:b/96 any**.

D)

Apply the following commands globally on SW1:

ipv6 mld vlan 76 immediate-leave

ipv6 mld snooping

E)

Apply the following commands globally on SW1:
ipv6 mld snooping multicast optimise-multicast-flood
ipv6 mld snooping fast-leave group security_access_list

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

ANSWER: A C

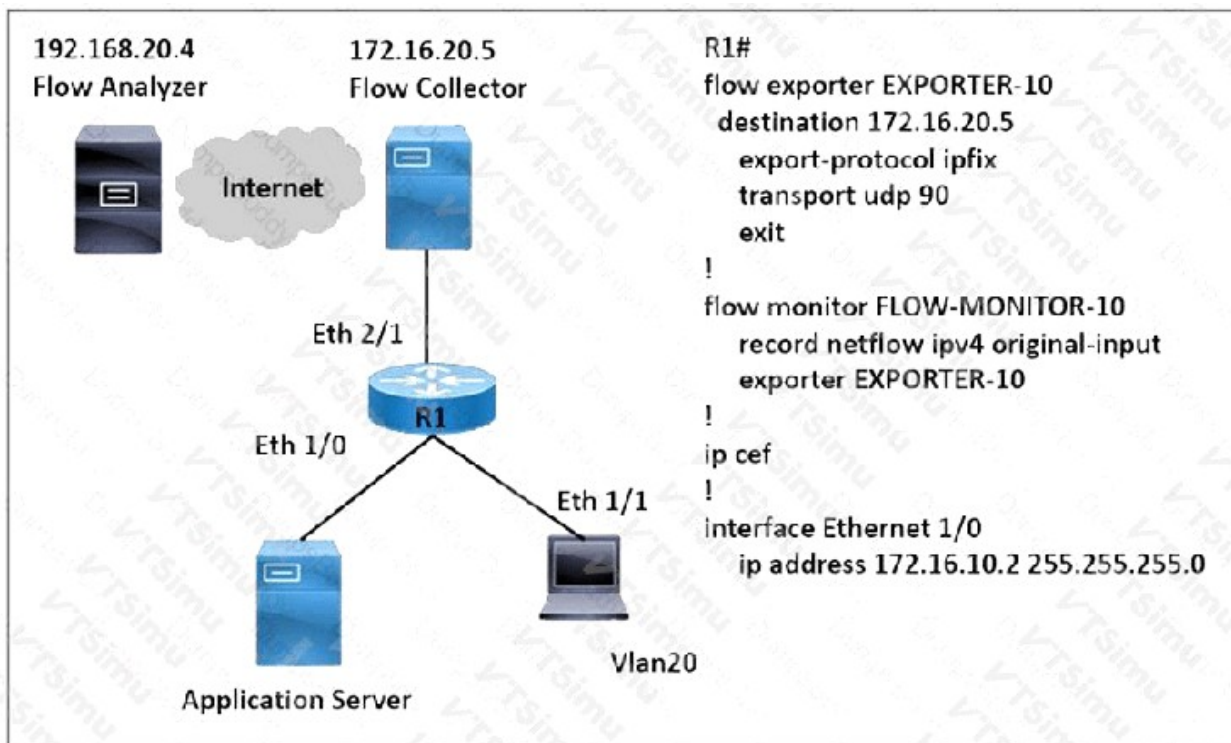
Explanation:

The configurations represented by the selected choices combine IPv6 Multicast Listener Discovery snooping for VLAN 76 with immediate-leave processing. MLD snooping examines MLD membership reports and maintains multicast forwarding state at Layer 2. Consequently, video multicast frames are forwarded only through switch interfaces with receivers that have joined the relevant IPv6 multicast group, rather than being flooded throughout VLAN 76. This directly reduces unnecessary bandwidth consumption, an especially important consideration for IPTV traffic.

Immediate leave complements snooping by allowing the switch to remove a port's multicast forwarding state as soon as it receives a leave indication for the group. The switch therefore stops sending the video stream to that port without waiting for the normal listener-query and timer process to expire. This provides the requested minimal leave latency when a subscriber stops watching or changes channels. Immediate leave is appropriate where a port represents a single receiver, or where the deployment ensures that removing the port immediately cannot disconnect another interested listener behind that port.

Cisco documents MLD snooping operation and VLAN-level configuration in its [IPv6 MLD Snooping Configuration Guide](#). The protocol behavior underlying listener reports and leave processing is defined by the IETF in [RFC 3810](#).

QUESTION NO: 10



Refer to the exhibit. A network engineer wants to monitor traffic from the application server and send the output to the external monitoring device at 172.16.20.5. Application server traffic should pass through the R1 Eth2/1 interface for further analysis after it is monitored. Which configuration must be applied on the R1 router?

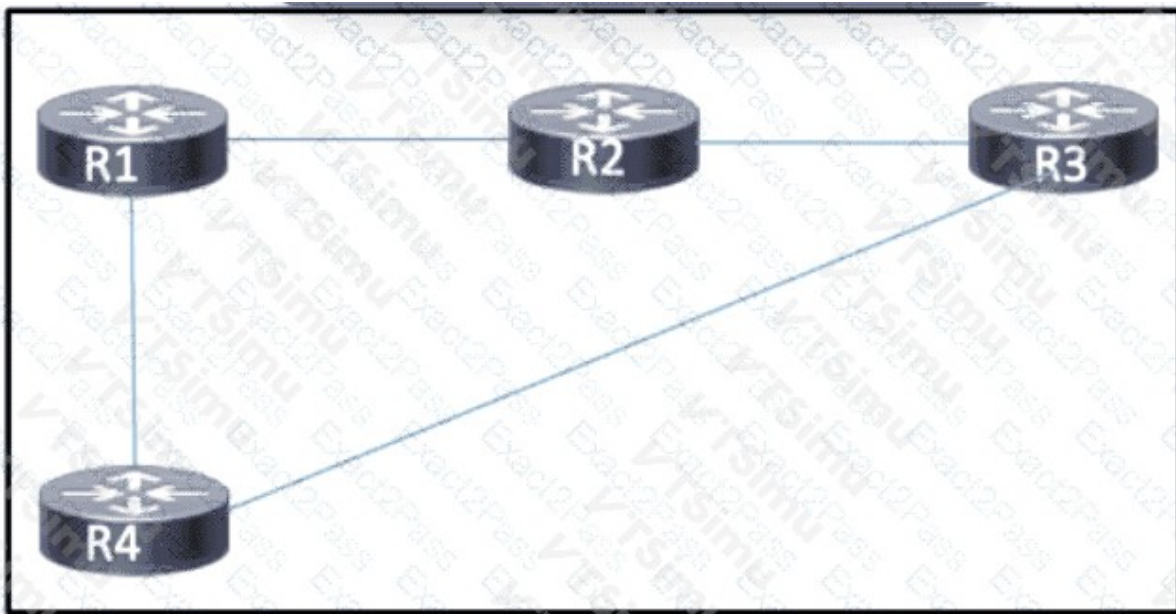
- A. Configure the FLOW-MONITOR-20 command.
- B. Configure the flow exporter EXPORTER-10 destination 192.168.20.4 command.
- C. Configure the ip flow monitor FLOW-MONITOR-10 input command on the Ethernet1/0 interface.
- D. Configure the ip flow monitor FLOW-MONITOR-10 output command on the Ethernet 2/1 interface.

ANSWER: C

Explanation:

Configure the `ip flow monitor FLOW-MONITOR-10 input` command on the Ethernet1/0 interface. is correct because Flexible NetFlow is applied in the direction in which the router first receives the traffic to be measured. The application server is connected on the Ethernet1/0 side of R1, so its packets enter R1 through that interface. Applying the named flow monitor in the input direction causes R1 to create and maintain flow records for the application server traffic as it arrives, while normal forwarding continues toward Ethernet2/1 for the required downstream analysis. The flow monitor ties together the flow record and the exporter configuration; the exporter sends the generated flow data to the configured external collector at 172.16.20.5. This is telemetry export, not packet redirection or traffic interception, so the application traffic can continue through Ethernet2/1 according to the routing table. Cisco's Flexible NetFlow documentation describes applying a flow monitor to an interface with the `ip flow monitor` command and selecting the `input` or `output` direction appropriate to the traffic observation point. See [Cisco Flexible NetFlow Overview](#) and [Cisco Flexible NetFlow Flow Monitor Configuration](#).

QUESTION NO: 11



Refer to the exhibit. MPLS is configured within the network. The routers use OSPF to exchange routing information. Hosts that are connected to routers R1 and R3 use routers R2 and R4 to access servers that run a variety of intranet applications. A network engineer must ensure that R1 and R3 maintain an LDP session between them to support the flow of traffic between hosts and servers in case of failure on the path between R1 and R3. Which task should the engineer perform to provide LDP bindings in case of failure on the path between R1 and R3?

- A. Implement session protection to send targeted hellos between R1 and R3.
- B. Implement multicast in the network to share data between hosts and servers without targeted routing.
- C. Implement LDP sync on the four devices to give MPLS and OSPF the ability to maintain continuous uptime.

D. Implement BFD to monitor and detect link failures between R1 and R3.

ANSWER: A

Explanation:

Implement session protection to send targeted hellos between R1 and R3. LDP session protection is designed to preserve an established LDP session when the directly connected link or its LDP discovery adjacency fails but an alternate IP path between the LDP peers remains available. When session protection is enabled, the routers use targeted LDP Hello messages sent to the peer's LDP router ID. Those targeted Hellos allow the LDP TCP session to remain established over the alternate routed path while OSPF reconverges around the failed path. As a result, R1 and R3 retain their learned label bindings rather than tearing down the LDP session and relearning labels after connectivity is restored. This minimizes MPLS forwarding disruption for traffic between the hosts and the intranet servers. The protection mechanism is particularly appropriate where resilient IGP connectivity exists between the two routers but the physical or directly adjacent LDP path can fail. Cisco documents LDP session protection as maintaining the LDP session through targeted Hellos following loss of the link Hello adjacency; see the [Cisco LDP Session Protection configuration guide](#) and the [Cisco MPLS LDP overview](#).

QUESTION NO: 12 - (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

Tasks

-

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

1. Establish R1 and R2 OSPF adjacency. All interfaces must be advertised in OSPF by using the OSPF interface command method. Use Loopback0 as the OSPF ID.
2. There must be no DR/BDR elections in OSPF Area 0 when establishing the neighbor relationship between R1 and R2. OSPF must not generate the host entries /32 for the adjacent interfaces.
3. Enable OSPF MD5 Authentication between both routers at the interface level with password C1sc0!

ANSWER: See the explanation for the answer

Explanation:

Configure OSPF directly under the interfaces (do not use `network` statements). Apply the configuration on both R1 and R2. Use the actual interface connected to the other router in place of `<R1-R2-link-interface>`, and use each router's configured Loopback0 address for its router ID.

R1

```
enable
configure terminal
router ospf 1
  router-id <R1-Loopback0-IP-address>
exit
interface Loopback0
  ip ospf 1 area 0
exit
interface <R1-R2-link-interface>
  ip ospf 1 area 0
  ip ospf network point-to-point
  ip ospf authentication message-digest
  ip ospf message-digest-key 1 md5 C1sc0!
end
copy running-config startup-config
```

R2

```
enable
configure terminal
router ospf 1
  router-id <R2-Loopback0-IP-address>
exit
interface Loopback0
  ip ospf 1 area 0
exit
interface <R1-R2-link-interface>
  ip ospf 1 area 0
  ip ospf network point-to-point
  ip ospf authentication message-digest
  ip ospf message-digest-key 1 md5 C1sc0!
end
copy running-config startup-config
```

If either router has any additional routed interfaces, advertise each one by entering the interface and applying:

```
interface <additional-interface>
  ip ospf 1 area 0
```

Verify the adjacency and OSPF interface type:

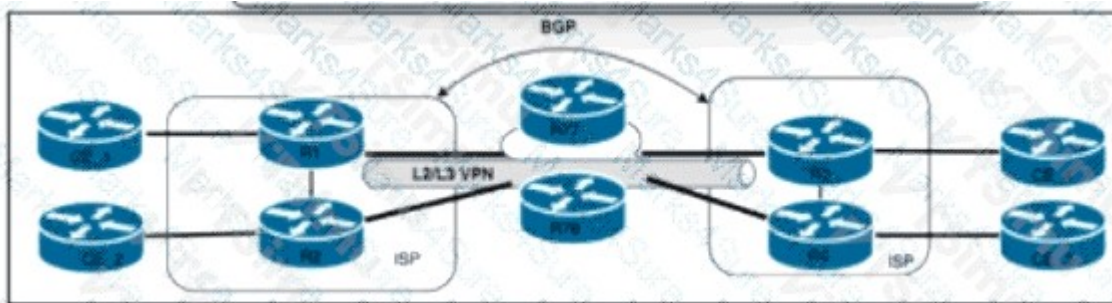
```
show ip ospf neighbor
show ip ospf interface <R1-R2-link-interface>
show ip route ospf
```

`ip ospf network point-to-point` prevents DR/BDR election on the R1-R2 Ethernet segment and prevents OSPF from advertising host (/32) entries for the adjacent link interfaces.

`ip ospf authentication message-digest` plus the identical MD5 key on both ends enables interface-level OSPF MD5 authentication.

The `router-id` must exactly match the IP address configured on that router's Loopback0.

QUESTION NO: 13



Refer to the exhibit. Company A's network architect is implementing the capabilities that are described in RFC 4379 to provide new options for troubleshooting the company's MPLS network. MPLS LSPs are signaled via the BGP routing protocol with IPv4 unicast prefix FECs. Due to company security policies, access is controlled on the control plane and management plane. Which two tasks must the architect perform to achieve the goal? (Choose two.)

- A. Open UDP port 3503.
- B. Open MCP port 3505.
- C. Implement OAM on all MPLS-enabled routers.
- D. Implement LSP Ping on PE routers only.
- E. Implement VCCV on core routers only.

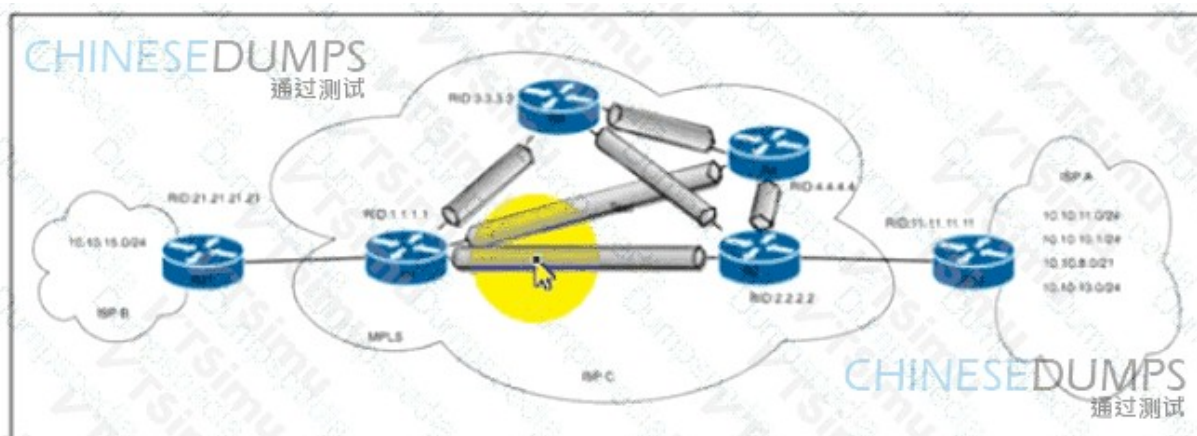
ANSWER: A D

Explanation:

Open UDP port 3503. RFC 4379 specifies MPLS LSP Ping and traceroute using MPLS Echo Request and Echo Reply messages encapsulated in UDP. UDP destination port 3503 is the well-known MPLS LSP Echo port. Because the network applies control-plane and management-plane filtering, the relevant ACLs and control-plane protection policies must permit this MPLS OAM traffic. This enables an ingress router to send an MPLS Echo Request toward the selected FEC and receive the resulting diagnostic response.

Implement LSP Ping on PE routers only. For BGP-signaled labeled-unicast LSPs using IPv4-prefix FECs, the ingress PE originates the test and the egress PE validates the FEC and produces the MPLS Echo Reply. The PE routers therefore need RFC 4379 LSP Ping capability for the BGP IPv4-prefix FEC endpoints. Transit P routers continue normal MPLS forwarding and TTL processing for the labeled packets; they do not need to originate or terminate the BGP IPv4-prefix FEC LSP Ping operation. RFC 4379 defines the echo protocol and its UDP transport, including the registered port usage. See [RFC 4379: Detecting Multi-Protocol Label Switched LSPs Data Plane Failures](#) and Cisco's [MPLS OAM configuration guide](#).

QUESTION NO: 14



Refer to the exhibit An engineer at ISP C is configuring a new interconnection with ISPs A and B using the BGP protocol After the initial configuration the engineer noticed high memory usage and an abnormally large LIB table on router R2 Which two actions must the engineer take on R2 to minimize memory usage? (Choose two.)

- A. Configure Extended ACL 101 with accepted prefixes.
- B. Configure the mpls ldp neighbor 11.11.11.11 labels accept 1 command.
- C. Configure Standard ACL 1 with accepted prefixes.
- D. Configure the mpls ldp neighbor 1.1.1.1 labels accept 101 command.
- E. Configure the mpls ldp neighbor 21.21.21.21 labels accept 101 command.

ANSWER: B C

Explanation:

Configuring Standard ACL 1 with the accepted prefixes and applying `mpls ldp neighbor 11.11.11.11 labels accept 1` restricts the label bindings that R2 retains in its Label Information Base (LIB). LDP normally receives label mappings for FECs from a neighbor and stores those mappings, which can consume substantial memory when a peer advertises bindings for a large number of prefixes. A standard IP access list used with LDP label acceptance identifies the FEC prefixes whose label mappings are useful to R2. ACL 1 should therefore permit only the required prefixes; those permitted FECs are the bindings that R2 accepts from the specified LDP peer. Applying the filter specifically to neighbor 11.11.11.11 makes the policy effective for the label advertisements received over that LDP session, reducing the number of LIB entries retained while preserving labels for the required reachable prefixes. Cisco documents LDP label filtering as a mechanism for controlling accepted label bindings and their associated resource use. See the [Cisco IOS MPLS LDP Configuration Guide](#) and the [Cisco IOS MPLS LDP Command Reference](#).

QUESTION NO: 15

Which module refers to the network automation using Ansible?

- A. the iosxr_system module to collect facts from remote devices
- B. the iosxr_user module to manage banners for users in the local database
- C. the iosxr_logging module to run debugging for severity levels 2 to 5
- D. the iosxr_command module to run commands on remote devices

ANSWER: D

Explanation:

The `iosxr_command` module is the appropriate Ansible module for issuing operational commands to Cisco IOS XR devices. In an automation playbook, it connects to the managed IOS XR router and runs one or more CLI commands, returning the command output in structured Ansible results. This is useful for repeatable operational tasks such as checking interface status, verifying routing information, collecting platform details, or validating the state of a service after a change. Commands can be supplied as a list, and the module can evaluate returned output against conditions through options such as `wait_for`, allowing a playbook to confirm that the desired network state has been reached. The module therefore supports agentless network automation: Ansible controls the device over its configured network connection rather than requiring an Ansible agent on the router. For current collection-based automation, Cisco IOS XR modules are documented in the `cisco.iosxr` Ansible collection. See the official [Ansible iosxr_command module documentation](#) and the [Ansible Cisco IOS XR platform guide](#).

QUESTION NO: 16

What is the primary role of a BR router in a 6rd environment?

- A. It provides connectivity between end devices and the IPv4 network.
- B. It embeds the IPv4 address in the 2002::/16 prefix.
- C. It connects the CE routers with the IPv6 network.
- D. It provides IPv4-in-IPv6 encapsulation

ANSWER: C

Explanation:

It connects the CE routers with the IPv6 network. In an IPv6 Rapid Deployment (6rd) design, the Border Relay (BR) is the provider-side gateway between the 6rd domain and the native IPv6 Internet or another native IPv6 network. Customer-edge (CE) routers use 6rd to carry IPv6 packets across an IPv4-only provider infrastructure by encapsulating the IPv6 packet in an IPv4 header. When traffic is destined beyond the 6rd domain, the BR receives the tunnel traffic, removes the IPv4 encapsulation, and forwards the original IPv6 packet using normal IPv6 routing. For traffic arriving from the native IPv6 side and destined for a 6rd customer prefix, the BR identifies the applicable 6rd delegation and sends the packet through the IPv4 network toward the relevant CE router. Thus, the BR supplies the essential interconnection point that lets 6rd CE routers reach native IPv6 destinations while the provider access/core transport remains IPv4. RFC 5969 defines the 6rd architecture and describes the BR's forwarding role: [RFC 5969](#). Cisco's implementation guidance is available in its [IPv6 6rd documentation](#).

QUESTION NO: 17

What are two characteristics of MPLS TE tunnels? (Choose two)

- A. They require EIGRP to be running in the core.
- B. They use RSVP to provide bandwidth for the tunnel.
- C. They are run over Ethernet cores only.
- D. The headend and tailend routes of the tunnel must have a BGP relationship
- E. They are unidirectional

ANSWER: B E

Explanation:

MPLS Traffic Engineering tunnels conventionally use RSVP with Traffic Engineering extensions (RSVP-TE) to signal the label-switched path and request resources along that path. The tunnel headend calculates or selects an explicit path subject to TE constraints, and RSVP-TE PATH and RESV signaling establish state at the participating routers. A configured bandwidth value can therefore be signaled as a reservation request, allowing the network to perform admission control based on available reservable bandwidth. This is the foundational control-plane behavior for RSVP-TE MPLS tunnels, as specified by [RFC 3209](#) and described in Cisco's [MPLS Traffic Engineering documentation](#).

An MPLS TE tunnel is also unidirectional: it represents a headend-to-tailend RSVP-TE LSP and carries traffic only in that signaled direction. When bidirectional traffic engineering is required, the reverse direction is implemented as a separate tunnel/LSP with its own headend, tailend, path selection, labels, and potentially separate bandwidth reservation. This directional model lets each direction follow an independently engineered path and use independently determined resource constraints.

QUESTION NO: 18

```

PE-A#show ip bgp vpnv4 vrf Customer-A neighbors 10.10.10.2 routes
BGP table version is 13148019, local router ID is 10.10.10.10
Status codes: s suppressed, d damped, h history, * valid, > best, i - internal,
               r RIB-failure, S Stale, m multipath, b backup-path, f RT-Filter,
               x best-external, a additional-path, c RIB-compressed,
Origin codes: i - IGP, e - EGP, ? - incomplete
RPKI validation codes: V valid, I invalid, N Not found

   Network          Next Hop          Metric LocPrf Weight Path
Route Distinguisher: 65000:1111 (default for vrf Customer-A)
*> 192.168.0.0/19    10.10.10.2          0         0 4282 65001 ?
*> 192.168.0.0/17    10.10.10.2          0         0 4282 65001 ?
*> 192.168.0.0/16    10.10.10.2          0         0 4282 65001 ?

Total number of prefixes 5

PE-A#config t
Enter configuration commands, one per line. End with CNTL/Z.
PE-A(config)#ip prefix-list ALLOW permit 192.168.0.0/16 ge 17 le 19
PE-A(config)#router bgp 65000
PE-A(config-router)#address-family ipv4 vrf Customer-A
PE-A(config-router-af)#neighbor 10.10.10.2 prefix-list ALLOW in

```

Refer to the exhibit. Which three outcomes occur if the prefix list is added to the neighbor? (Choose three.)

- A. 192.168.0.0/16 is denied.
- B. 192.168.0.0/16 is permitted.
- C. 192.168.0.0/19 is permitted
- D. 192.168.0.0/19 is denied.
- E. 192.168.0.0/17 is permitted
- F. 192.168.0.0/17 is denied.

ANSWER: A C F

Explanation:

The outcomes are **192.168.0.0/16 is denied**, **192.168.0.0/19 is permitted**, and **192.168.0.0/17 is denied**. The configured prefix-list entry uses 192.168.0.0/16 as its base prefix and permits only matching prefixes whose prefix length meets the configured length constraint. With a permitted length of 19, the 192.168.0.0/19 route matches both the base network and the required mask length, so it is accepted when the policy is applied to the BGP neighbor.

A prefix list does not treat the base prefix as an automatic permit when a `ge` or `le` condition narrows the acceptable prefix-length range. Consequently, the exact /16 route and the /17 route do not match the permitted /19 length. Cisco prefix lists include an implicit deny for prefixes that do not match a permit statement, producing denial for those routes. When attached to the neighbor in the direction shown in the exhibit, this controls which matching BGP NLRI is accepted or advertised in that policy direction. Cisco documents prefix-list matching and the use of `ge` and `le` qualifiers in its [IP Prefix List configuration guide](#) and [BGP prefix filtering documentation](#).

QUESTION NO: 19

Refer to the exhibit.

```
R1
ip multicast-routing
ip pim rp-candidate GigabitEthernet1/0/0

interface g1/0/0
 ip pim sparse-mode

R2
ip multicast-routing
ip pim bsr-candidate GigabitEthernet1/0/0

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

An engineer configured multicast routing on client's network. What is the effect of this multicast implementation?

- A. R2 floods information about R1 throughout the multicast domain.
- B. R2 is unable to share information because the ip pim autorp listener command is missing.
- C. R1 floods information about R2 throughout the multicast domain.
- D. R2 is elected as the RP for this domain.

ANSWER: D

Explanation:

R2 is elected as the RP for this domain. Cisco Auto-RP uses Candidate-RP (C-RP) announcements and RP-Discovery messages to distribute rendezvous-point information dynamically. A router configured as a candidate RP advertises its availability for the configured multicast group range. The Auto-RP mapping agent receives these candidate-RP announcements, selects the appropriate RP, and floods RP-Discovery messages through the multicast domain so multicast routers can learn the selected RP-to-group mapping.

In the depicted implementation, the candidate-RP and mapping-agent behavior results in R2 being selected as the rendezvous point. The RP is a fundamental component for Protocol Independent Multicast sparse mode because receivers initially join toward it and sources register multicast traffic with it. Routers that receive the Auto-RP discovery information can therefore build the PIM shared tree toward R2 for groups in the advertised range. Cisco documents the Auto-RP process and the candidate-RP/mapping-agent roles in its [Auto-RP configuration guide](#) and describes RP operation in the [PIM rendezvous point overview](#).

QUESTION NO: 20

Which feature describes the weight parameter for BGP path selection?

- A. Its value is local to the router

- B. Its value is set either locally or globally.
- C. Its default value is 0.
- D. Its value is global to the router.

ANSWER: A

Explanation:

Its value is local to the router is correct because BGP weight is a Cisco-specific, nontransitive path-selection parameter that affects only the Cisco router where it is configured. When that router has multiple BGP paths for the same destination, it prefers the path with the highest weight before evaluating standard BGP attributes such as local preference, AS-path length, origin, and MED. Weight is not included in BGP update messages and is therefore neither propagated to BGP peers nor used by other routers when they make their own best-path decisions. This makes weight useful when an administrator needs to influence outbound path selection on one router without changing routing policy across the autonomous system or communicating the preference externally. Cisco documents weight as a locally significant Cisco attribute and places it first in its BGP best-path selection process. See Cisco's [BGP Best Path Selection Algorithm](#) and the [Cisco IOS XE BGP Path Selection documentation](#).

QUESTION NO: 21

What is a characteristics of the Pipe model for MPLS QoS?

- A. The same QoS policy is applied to all customer traffic on the egress PE.
- B. If the outer EXP is changed, it is copied to the DSCP value.
- C. The MPLS EXP bits are set by the CE.
- D. The DSCP value determines how the packet is forwarded
- E. The ingress PE sets the MPLS EXP bits according to the provider QoS policy, independently of the customer IP DSCP value.

ANSWER: E

Explanation:

In the MPLS DiffServ Pipe model, the provider treats the MPLS network as a QoS "pipe" whose forwarding behavior is controlled independently of the customer packet's original IP DSCP value. At the ingress PE, the provider applies its configured QoS policy and sets the MPLS Traffic Class field (historically called EXP bits) on the imposed MPLS label. MPLS routers in the provider core use that Traffic Class value to select the configured per-hop behavior, queue, and drop treatment. The customer IP header marking is preserved rather than being rewritten from the MPLS Traffic Class value when the label is removed. This separation lets a service provider deliver an MPLS transport class while retaining the customer's IP marking in the encapsulated packet. Cisco documents MPLS DiffServ tunneling modes and the use of MPLS EXP/Traffic Class marking at label imposition in its [MPLS DiffServ Tunneling Modes documentation](#). The Pipe model is also defined by the IETF as a model in which the MPLS tunnel's DiffServ treatment is independent of the encapsulated IP packet's DS field; see [RFC 3270](#).

QUESTION NO: 22

Which two statements about MPLS Layer 3 VPN route targets are true? (Choose two.)

- A. They are BGP extended communities used to control VPN route distribution.
- B. A VRF imports a VPN route when the route carries a matching import route target.
- C. They make overlapping customer IPv4 prefixes unique in the VPNv4 address family.
- D. They are distributed by LDP with transport labels.

E. They must be globally unique for every VRF in the provider network.

ANSWER: A B

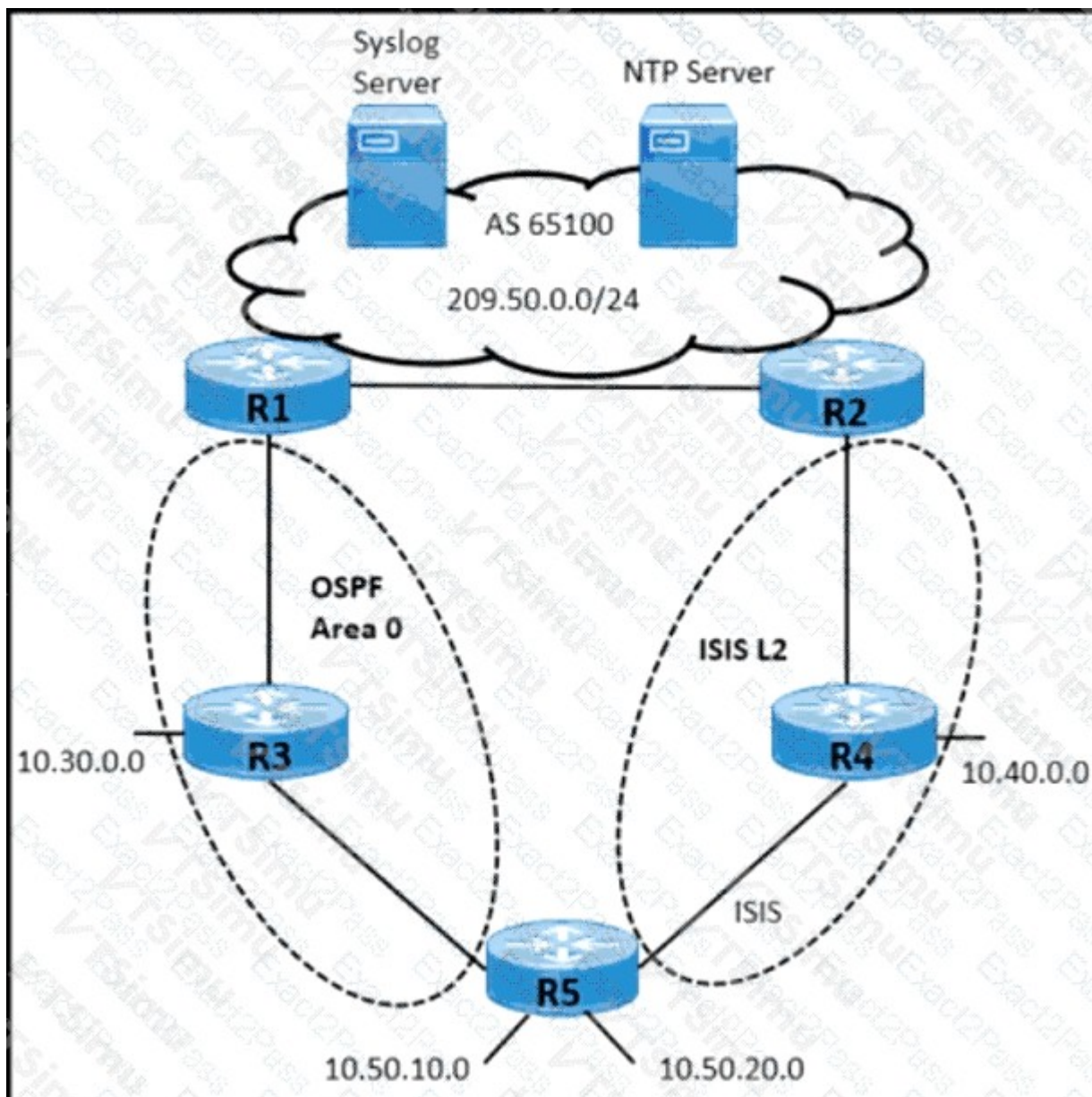
Explanation:

Route targets are BGP extended communities used to control the import and export of VPN routes between VRFs. A PE attaches export route targets to VPNv4 or VPNv6 routes when it advertises them through multiprotocol BGP. A receiving PE imports a VPN route into a local VRF only when the route carries a route target that matches an import route target configured for that VRF.

This mechanism supports controlled connectivity among customer VRFs and shared-service VRFs. Route targets are different from route distinguishers: a route distinguisher makes an IPv4 or IPv6 prefix unique in the VPN address family, whereas route targets define VPN route distribution policy. Cisco describes route-target import and export policy in the [Cisco IOS XE MPLS Layer 3 VPN VRF Configuration Guide](#).

QUESTION NO: 23

Refer to the exhibit.



A network operator working for a telecommunication company with an employee ID: 4350:47:853 must implement an IGP solution based on these requirements:

- Subnet 10.50.10.0 traffic must exit through the R1 router to connect with the Syslog server.
- Subnet 10.50.20.0 traffic must exit through the R2 router to connect with the NTP server.
- In case of link failure between R2 and R4, traffic must be routed via R1 and R3.

Which two configurations must be implemented on R5 to meet these requirements? (Choose two.)

- A. Apply a route policy to redistribute 10.50.0.0 prefixes in OSPF to ISIS and ISIS to OSPF.
- B. Apply a route policy to redistribute 10.50.20.0 from ISIS-L2 to OSPF Area 0 at a higher cost.
- C. Enable a route policy to advertise 10.50.20.0 in ISIS-L2 at a higher cost.
- D. Apply a route policy to redistribute 10.50.10.0 from OSPF Area 0 to ISIS-L2 at a lower cost.
- E. Enable a route policy to advertise 10.50.10.0 In OSPF Area 0 at a low cost.

ANSWER: C E

Explanation:

R5 must control the metric with which each destination prefix is advertised into the IGP domain that supplies the relevant exit path. Enabling a route policy to advertise 10.50.10.0 In OSPF Area 0 at a low cost makes the path toward the Syslog subnet through R1 the preferred OSPF route. OSPF calculates the shortest path by accumulating interface and advertised costs, so a deliberately low advertised metric ensures that this intended exit is selected.

Enabling a route policy to advertise 10.50.20.0 in ISIS-L2 at a higher cost provides the required IS-IS Level-2 metric treatment for the NTP destination while preserving an alternate reachable path. When the R2-to-R4 link is unavailable, normal SPF recalculation selects the remaining topology path through R1 and R3. This design uses route-policy-controlled metrics to express the intended primary exit behavior and allows IGP convergence to provide the specified failure path without requiring static routing.

Cisco documents IS-IS metric and route-policy controls in its [IS-IS configuration guide](#). OSPF shortest-path metric behavior is defined in [RFC 2328](#).

QUESTION NO: 24 - (DRAG DROP)

DRAG DROP

Drag and drop the LDP features from the left onto the correct usages on the right.

Select and Place:

Answer Area

session protection	It prevents valid routes from being overwritten with new ones until labels are assigned.
IGP synchronization	It allows stale label bindings to be used for a period of time while an LDP neighbor is unreachable.
targeted-hello accept	It uses LDP Targeted hellos to protect LDP sessions.
graceful restart	It uses LDP to form neighborhood between non-directly connected routers.

ANSWER:

Answer Area

session protection	IGP synchronization
IGP synchronization	graceful restart
targeted-hello accept	session protection
graceful restart	targeted-hello accept

Explanation:

LDP IGP synchronization ensures that an IGP does not begin using a newly available path for MPLS traffic before the required LDP label information is ready. In practical terms, this holds back the relevant IGP route advertisement or metric change until label bindings have been established, preventing a valid route from being preferred before packets can be label-switched correctly. This avoids transient forwarding failures during link startup or topology changes.

LDP graceful restart preserves forwarding continuity when an LDP peer becomes temporarily unreachable or restarts. The router can retain previously learned, stable label bindings for a configured recovery period, allowing labeled forwarding to continue while the LDP session is restored. Session protection uses targeted LDP Hellos to maintain knowledge of a peer and facilitate rapid session reestablishment when the directly connected adjacency is disrupted. A targeted Hello is also the mechanism used to establish LDP peer relationships between routers that are not directly connected; targeted-hello acceptance permits those incoming targeted Hello messages to be accepted. Cisco's LDP documentation describes targeted discovery and session protection behavior in its [LDP Session Protection configuration guide](#), while the interaction of LDP and IGP route advertisement is covered in the [LDP IGP Synchronization guide](#).

QUESTION NO: 25

Refer to the exhibit:

```
class-map match-any class1
match protocol ipv4
match qos-group 4
```

A network engineer is implementing QoS services. Which two statements about the QoS-group keyword on Cisco IOS XR are true? (Choose two)

- A. The QoS group numbering corresponds to priority level
- B. QoS group marking occurs on the ingress
- C. It marks packets for end to end QoS policy enforcement across the network
- D. QoS group can be used in fabric QoS policy as a match criteria
- E. It cannot be used with priority traffic class

ANSWER: B D

Explanation:

“QoS group marking occurs on the ingress” is correct because a QoS group is an internally significant packet classification value assigned as traffic enters the Cisco IOS XR router. The value is carried in the router’s internal packet metadata, allowing subsequent QoS processing stages to identify traffic according to the ingress policy decision. It is therefore useful when the original packet marking alone is not the desired basis for internal classification.

“QoS group can be used in fabric QoS policy as a match criteria” is also correct. After an ingress policy sets the internal QoS-group value, a fabric QoS policy can match that value and apply the appropriate internal forwarding, queuing, congestion-management, or scheduling treatment. This supports a modular QoS design in which ingress classification and marking establish the internal traffic identity, while fabric policy enforces the associated handling through the router. Cisco IOS XR QoS documentation describes QoS groups as internal markings that can be set and matched by policy maps, including for fabric QoS workflows. See Cisco’s [IOS XR QoS Configuration Guide](#) and the [policy-map configuration documentation](#).

QUESTION NO: 26



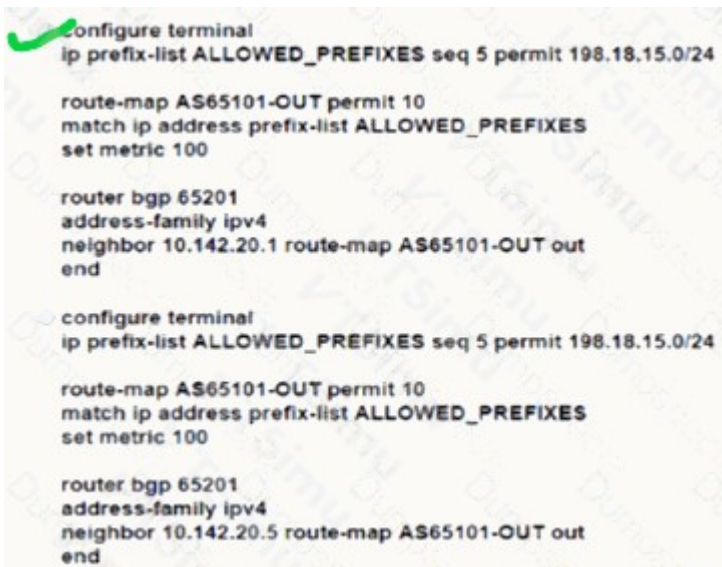
Refer to the exhibit an engineer working for a private telecommunication company with an employee Id: 4065:96:080 upgrades the WAN link between routers ASBR-101 and ASBR-201 to 1Gb by Installing a new physical connection between the Gi3 Interfaces. Which BGP attribute must the engineer configure on ASBR-201 so that the existing WAN link on Gi2 is maintained as a backup?

- configure terminal
ip prefix-list ALLOWED_PREFIXES seq 5 permit 198.18.15.0/24

route-map AS65101-OUT permit 10
match ip address prefix-list ALLOWED_PREFIXES
set as-path prepend 65101 65101

router bgp 65201
address-family ipv4
neighbor 10.142.20.1 route-map AS65101-OUT out
end
- configure terminal
ip prefix-list ALLOWED_PREFIXES seq 5 permit 198.18.15.0/24

route-map AS65101-OUT permit 10
match ip address prefix-list ALLOWED_PREFIXES
set as-path prepend 65101 65101



```

configure terminal
ip prefix-list ALLOWED_PREFIXES seq 5 permit 198.18.15.0/24

route-map AS65101-OUT permit 10
match ip address prefix-list ALLOWED_PREFIXES
set metric 100

router bgp 65201
address-family ipv4
neighbor 10.142.20.1 route-map AS65101-OUT out
end

configure terminal
ip prefix-list ALLOWED_PREFIXES seq 5 permit 198.18.15.0/24

route-map AS65101-OUT permit 10
match ip address prefix-list ALLOWED_PREFIXES
set metric 100

router bgp 65201
address-family ipv4
neighbor 10.142.20.5 route-map AS65101-OUT out
end

```

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

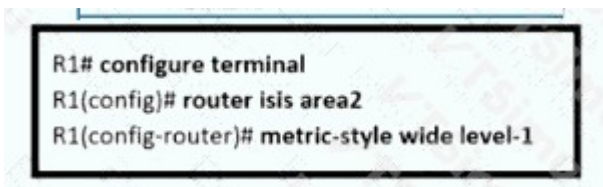
ANSWER: C

Explanation:

Weight is the appropriate BGP attribute because the path-selection requirement is local to ASBR-201: it must prefer the newly installed Gi3 path while retaining the Gi2 peering and its routes as an available standby path. Cisco BGP evaluates Weight before the other standard best-path attributes. Therefore, ASBR-201 can assign a higher Weight to routes learned from the Gi3 neighbor, causing those routes to be selected as the active best paths. The same prefixes learned through Gi2 remain in the BGP table but are not installed as the preferred forwarding paths while Gi3 is available. If Gi3 or its BGP session fails, the Gi2-learned paths can become best and restore connectivity over the existing WAN link. Weight is a Cisco-specific attribute and is significant only on the router where it is configured; this makes it well suited for controlling ASBR-201's own exit-path decision without requiring a policy change to be propagated elsewhere. Cisco documents Weight as the first BGP best-path criterion and notes that a higher value is preferred. See Cisco's [BGP Best Path Selection Algorithm](#) and the [Cisco IOS XE BGP Configuration Guide](#).

QUESTION NO: 27

Refer to the exhibit.



```

R1# configure terminal
R1(config)# router isis area2
R1(config-router)# metric-style wide level-1

```

An engineer is configuring multiprotocol IS-IS for IPv6 on router R1. Which additional configuration must be applied to the router to complete the task?

```

R1# configure terminal
R1(config)# router isis area1
R1(config-router)# metric-style wide level-1
R1(config-router)# address-family ipv6
R1(config-router-af)# multi topology

R1# configure terminal
R1(config)# router isis area2
R1(config-router)# metric-style wide
R1(config-router)# address-family ipv6
R1(config-router-af)# multi topology

R1# configure terminal
R1(config)# router isis area1
R1(config-router)# metric-style wide level-2
R1(config-router)# address-family ipv6
R1(config-router-af)# multi-topology

R1# configure terminal
R1(config)# router isis area2
R1(config-router)# address-family ipv6
R1(config-router-af)# multi-topology

```

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

ANSWER: D

Explanation:

Multitopology IS-IS requires the IPv6 unicast address family to be enabled under the IS-IS routing process and configured to use a separate topology with the `multi-topology` command. This causes IS-IS to advertise and calculate IPv6 reachability independently from the base IPv4 topology, while retaining a single IS-IS protocol instance. The configuration is completed in the IPv6 address-family context of the IS-IS process, for example: `router isis <tag>`, followed by `address-family ipv6 unicast` and `multi-topology`. IPv6 must also be enabled for IS-IS on the participating interfaces so that IPv6 link-state information and IPv6 prefixes can be advertised through the configured IPv6 topology. This design allows IPv4 and IPv6 to use different SPF calculations and potentially different paths, metrics, or link participation where required. Cisco documents the IS-IS IPv6 unicast address family and its multitopology capability in the IOS XE IS-IS command reference and configuration guide. See the [Cisco IS-IS command reference](#) and the [Cisco IS-IS IPv6 configuration guide](#).

QUESTION NO: 28

Which two statements about gRPC Network Management Interface are true? (Choose two.)

- A. It supports streaming telemetry through the Subscribe operation.
- B. It can modify configuration through the Set operation.
- C. It requires device configuration to be collected by CLI screen scraping.
- D. It uses SNMP traps as its transport protocol.
- E. It cannot retrieve operational data from a device.

ANSWER: A B

Explanation:

gRPC Network Management Interface, commonly called gNMI, is a gRPC-based protocol used for model-driven configuration management and telemetry. A gNMI client can retrieve data with the Get operation and can modify configuration with the Set operation. Data is identified by structured paths that are commonly based on YANG-modeled schema.

The Subscribe operation supports streaming telemetry from a device to a collector. A subscription can deliver periodic samples or updates when data changes, depending on the configured subscription mode and device support. gNMI does not depend on CLI screen scraping, and it is distinct from SNMP even though both can provide management visibility. See the [gNMI specification](#) and [Cisco IOS XR gNMI telemetry documentation](#).

QUESTION NO: 29

```
RP/0/RSP0/CPU0:JFK-PE#show mpls ldp bindings 192.168.10.10/32
Fri Nov 11 21:02:33.124 UTC
192.168.10.10/32, rev 2
  Local bindings: label: ImpNull
  Remote bindings: (2 peers)
    Peer          Label
    -----
    10.10.10.2:0   562656
    10.10.10.5:0   378337
```

Refer to the exhibit. After implementing a new design for the network, a technician reviews the pictures CLI output as part of the MOP. Which two elements describe what the technician can ascertain from the ImpNull output? (Choose two.)

- A. Ultimate Hop Popping is in use for the prefix displayed.
- B. Penultimate Hop Popping is in use for the prefix displayed.
- C. Label 0 is used for the prefix displayed, but will not be part of the MPLS label stack for packets destined for 192.168.10.10.
- D. Label 3 is used for the prefix displayed, but will not be part of the MPLS label stack for packets destined for 192.168.10.10.
- E. Label 0 is used for the prefix displayed and will be part of the MPLS label stack for packets destined for 192.168.10.10.

ANSWER: B D

Explanation:

ImpNull denotes the MPLS *Implicit Null* label binding. Implicit Null has the reserved label value of 3 and is used to request that the upstream, penultimate router remove the top MPLS label before forwarding the packet to the egress router. This behavior is Penultimate Hop Popping. It reduces processing at the egress label-switching router because the egress receives a native IP packet rather than one carrying the transport label. Although Implicit Null is assigned reserved value 3 in MPLS signaling and label bindings, it is not encoded into the MPLS label stack on forwarded packets. Instead, its operational instruction is to pop the existing top label. Therefore, the CLI indication establishes both the use of Penultimate Hop Popping for the displayed prefix and the use of the reserved Implicit Null label value without that label appearing in the packet's transmitted MPLS stack. The IETF MPLS label-stack specification defines label value 3 as Implicit NULL and states that it is not actually encoded in the stack: [RFC 3032](#). Cisco also documents Implicit Null as the mechanism that causes the penultimate router to pop the label: [Cisco MPLS label allocation documentation](#).

QUESTION NO: 30



```

CPE-1#show run int gig 0/0
interface GigabitEthernet0/0
 ip address 100.65.15.2 255.255.255.252
 negotiation auto
 ipv6 address 2001:DB8:0:A000:100:65:15:2/126
 service-policy output WAN-OUTPUT
end

CPE-1#show run int gig 0/1
interface GigabitEthernet0/1
 ip address 192.168.2.1 255.255.255.0
 negotiation auto
 ipv6 address 2001:DB8:0:A001:192:168:2:1/120
 service-policy input LAN-INPUT
end

CPE-1#show access-list
Standard IP access list SELF_V4
 10 permit 100.65.15.2
IPv6 access list SELF_V6
 permit ipv6 host 2001 :DB8:0:A000:100:65:15:2 any sequence 10

CPE-1#show policy map
Policy Map WAN OUTPUT

Policy Map LAN INPUT
  
```

Refer to the exhibit. A network engineer configures CPE-1 for QoS with these requirements:

IPv4 and IPv6 traffic originated by the CPE-1 WAN IP address must be marked with DSCP CS3.

IPv4 LAN traffic must be marked with DSCP CS1.

IPv6 LAN traffic must be marked with DSCP default.

Which configuration must the engineer implement on CPE-1?

- A.** class-map match-any SELF_TRAFFICmatch access-group name SELF_V4match access-group name SELF_V6class-map match-all V4_ TRAFFICmatch protocol ipclass-map match-all V6_ TRAFFICmatch protocol ipv6class-map match-all QG_4match qos-group 4class-map match-all QG_6match qos-group 6!policy-map LAN-INPUTclass V4_TRAFFICset qos-group 4class V
- B.** class-map match-all SELF_TRAFFICmatch access-group name SELF_V4match access-group name SELF_V6class-map match-all V4_ TRAFFICmatch protocol ipclass-map match-all V6_ TRAFFICmatch protocol ipv6class-map match-all QG_4match qos-group 4class-map match-all QG_6match qos-group 6!policy-map LAN-INPUTclass V4_TRAFFICset qos-group 4class V
- C.** class-map match-all SELF_TRAFFICmatch access-group name SELF_V4match access-group name SELF_V6class-map match-all V4_ TRAFFICmatch protocol ipclass-map match-all V6_ TRAFFICmatch protocol ipv6class-map match-all QG_4match qos-group 4class-map match-all QG_6match qos-group 6!policy-map LAN-INPUTclass V4_TRAFFICset qos-group 4class V
- D.** class-map match-any SELF_TRAFFICmatch access-group name SELF_V4match access-group name SELF_V6class-map match-all V4_ TRAFFICmatch protocol ipclass-map match-all V6_ TRAFFICmatch protocol ipv6class-map match-all

QG_4match qos-group 4class-map match-all QG_6match qos-group 6!policy-map LAN-INPUTclass V4_TRAFFICset qos-group 4class V

ANSWER: A

Explanation:

The configuration beginning with `class-map match-any SELF_TRAFFIC` correctly creates a single class for router-originated WAN traffic across both IP versions. The IPv4 and IPv6 ACL matches are alternatives: a packet sourced from the CPE WAN IPv4 address matches the IPv4 ACL, while a packet sourced from the CPE WAN IPv6 address matches the IPv6 ACL. Therefore, `match-any` is required for the combined self-traffic class. The corresponding egress policy can mark that class as DSCP CS3.

The design also separates LAN-originated IPv4 and IPv6 packets at ingress by protocol and assigns internal QoS groups. A subsequent WAN egress policy can match QoS group 4 and set DSCP CS1 for IPv4 LAN traffic, while matching QoS group 6 and setting DSCP default for IPv6 LAN traffic. This two-stage approach preserves the required distinction between traffic sourced by CPE-1 itself and traffic forwarded from the LAN, then applies the required DSCP values at the WAN egress point. Cisco MQC class maps support multiple match statements with `match-any` semantics for this type of alternative classification. See Cisco's [QoS classification documentation](#) and [QoS marking documentation](#).

QUESTION NO: 31 - (SIMULATION)



Guidelines

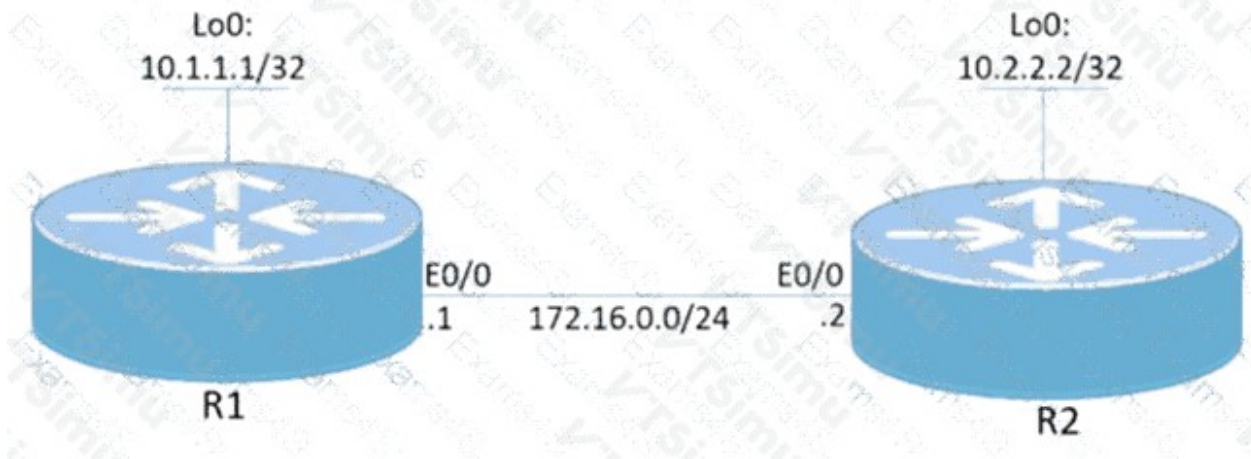
-

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

1. Establish R1 and R2 OSPF adjacency. All interfaces must be advertised in OSPF by using the OSPF interface command method. Use Loopback0 as the OSPF I

D.

2. There must be no DR/BDR elections in OSPF Area 0 when establishing the neighbor relationship between R1 and R2. OSPF must not generate the host entries /32 for the adjacent interfaces.

3. Enable OSPF MD5 Authentication between both routers at the interface level with password C1sc0!.

ANSWER: See the explanation for the answer

Explanation:

Configure OSPF process 10 directly under each required interface (do not use OSPF `network` statements). Set Ethernet to OSPF point-to-point so that no DR/BDR election occurs and OSPF advertises the transit network rather than creating /32 host routes for the Ethernet endpoints.

R1

```
enable
configure terminal
router ospf 10
  router-id 10.1.1.1
exit
interface loopback0
  ip ospf 10 area 0
exit
interface ethernet0/0
  ip ospf 10 area 0
  ip ospf network point-to-point
  ip ospf authentication message-digest
  ip ospf message-digest-key 1 md5 C1sc0!
end
write memory
```

R2

```
enable
configure terminal
router ospf 10
  router-id 10.2.2.2
exit
interface loopback0
  ip ospf 10 area 0
exit
interface ethernet0/0
  ip ospf 10 area 0
  ip ospf network point-to-point
  ip ospf authentication message-digest
  ip ospf message-digest-key 1 md5 C1sc0!
end
write memory
```

If OSPF was already running before the router IDs were configured, restart the process on each router so the new IDs take effect:

```
clear ip ospf process
```

Verify the result with:

```
show ip ospf neighbor
show ip ospf interface ethernet0/0
show ip route ospf
```

The neighbor state should be FULL/-; the dash confirms there is no DR or BDR role.

show ip ospf interface ethernet0/0 should show network type POINT_TO_POINT and message-digest authentication.

The OSPF routes should include the remote loopback /32 and the 172.16.0.0/24 transit subnet, not OSPF-generated /32 entries for 172.16.0.1 and 172.16.0.2.

QUESTION NO: 32

```
class-map match-any class1
match-protocol ipv4
match qos-group 4
```

Refer to the exhibit. A network engineer is implementing QoS services. Which two statements about the qos-group keyword on Cisco IOS XR are true?

(Choose two.)

- A. It marks packets for end-to-end QoS policy enforcement across the network.
- B. QoS group marking occurs on the ingress.
- C. The QoS group numbering corresponds to priority level.
- D. QoS group can be used in fabric QoS policy as match criteria.
- E. It cannot be used with priority traffic class.

ANSWER: B D

Explanation:

QoS group marking occurs on the ingress because qos-group is an internally significant QoS classification value assigned as packets enter the router. An ingress policy can classify traffic using packet attributes such as DSCP, precedence, MPLS EXP, or access-control criteria, and then set a QoS-group value. This preserves the classification through the IOS XR

forwarding pipeline without requiring a change to customer-visible packet markings. The value is therefore useful for carrying the result of ingress classification to later QoS processing stages within the device.

QoS group can be used in fabric QoS policy as match criteria because IOS XR fabric QoS uses this internal marking to identify traffic classes while packets traverse the router fabric and are handled at egress. A fabric policy can match the QoS-group value established on ingress and apply the appropriate traffic-class treatment, such as mapping to the desired internal traffic class and queueing behavior. This separation between ingress classification and fabric/egress treatment supports consistent QoS handling across distributed IOS XR hardware. Cisco's IOS XR QoS documentation describes QoS-group as internal marking used by QoS policies and fabric processing; see the [Cisco IOS XR QoS Configuration Guide](#) and the [classification and marking guidance](#).

QUESTION NO: 33 - (DRAG DROP)

DRAG DROP

Drag and drop the OSs from the left onto the correct descriptions on the right.

Select and Place:

Answer Area

IOS XR	It is a monolithic architecture that runs all modules on one memory space.
IOS	It runs over a Linux platform and pulls the system functions out of the main kernel and into separate processes.
IOS XE	It segments ancillary processes into separate memory spaces to prevent system crashes from errant bugs.

ANSWER:

Answer Area

IOS XR	IOS
IOS	IOS XE
IOS XE	IOS XR

Explanation:

Cisco IOS is correctly associated with the monolithic architecture description. Traditional IOS executes its operating-system functions and protocol components in a shared memory space, which is the defining model of a monolithic network operating system. Cisco IOS XE is correctly identified as the Linux-based evolution of IOS. IOS XE uses a Linux kernel and runs the

IOS control-plane software as a separate IOSd process, separating major system functions from the underlying kernel environment. This design supports improved software modularity while retaining familiar IOS operational behavior. Cisco IOS XR is correctly associated with segmentation of processes into separate memory spaces. IOS XR was designed for carrier-class routing and high availability, using a modular architecture with protected, isolated processes. That separation limits the impact of a faulty ancillary process, allowing the rest of the system and essential routing functions to continue operating more reliably. These architectural characteristics are especially important in service-provider networks, where process isolation, fault containment, and nonstop operation are fundamental requirements. Cisco describes IOS XR as a modular operating system designed for high availability and resiliency in its [IOS XR Software overview](#). Cisco also documents the Linux-based IOS XE software architecture and IOSd process model in the [Cisco IOS XE configuration documentation](#).

QUESTION NO: 34

```
R1
interface fastethernet1/0
  ip address 192.168.1.3 255.255.255.0
router bgp 65000
  router-id 192.168.1.1
  neighbor 192.168.1.2 remote-as 65012

R2
interface fastethernet1/0
  ip address 192.168.1.2 255.255.255.0
router bgp 65012
  router-id 192.168.1.1
  neighbor 192.168.1.3 remote-as 65000
  neighbor 192.168.1.3 local-as 65112
```

Refer to the exhibit. Assume all other configurations are correct and the network is otherwise operating normally.

Which conclusion can you draw about the neighbor relationship between routers R1 and R2?

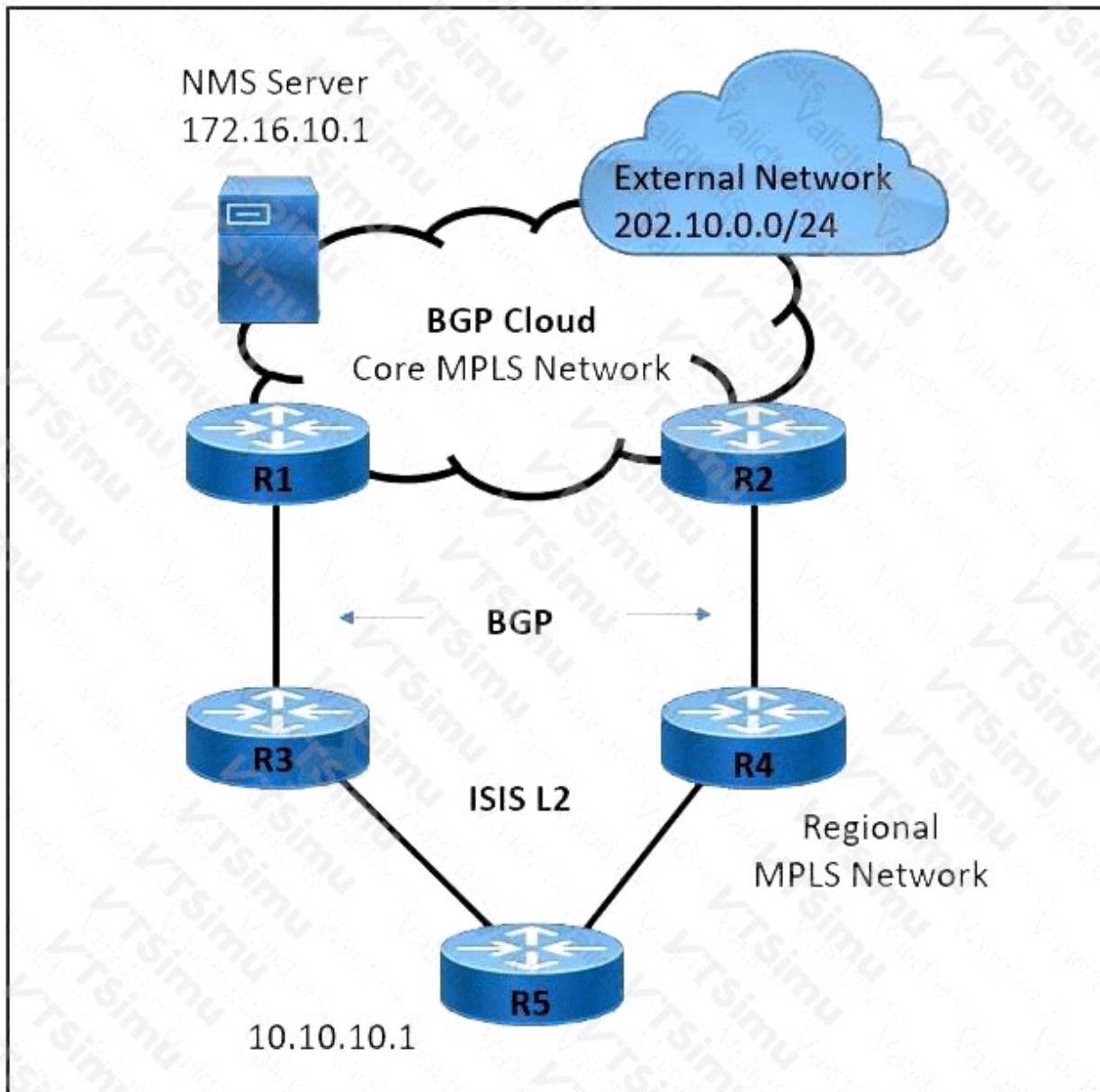
- A. The neighbor relationship is up.
- B. The neighbor relationship will be up only if the two devices have activated the correct neighbor relationships under the IPv4 address family.
- C. The neighbor is down because the local-as value for R2 is missing in the R1 neighbor statement.
- D. The neighbor relationship is down because R1 believes R2 is in AS 65012.

ANSWER: D

Explanation:

The neighbor relationship is down because R1 believes R2 is in AS 65012. In BGP, the autonomous system configured with the `remote-as` parameter identifies the AS that the local router expects the peer to advertise during BGP session establishment. R1 therefore expects the BGP OPEN message received from R2 to identify AS 65012. However, the configuration shown for R2 identifies a different BGP AS. This causes an AS-number mismatch during OPEN-message processing, so the BGP finite-state machine cannot establish the peering session.

The expected remote AS is a fundamental part of BGP neighbor configuration, both for external BGP peers and for internal BGP peers. The configured value must match the AS that the remote router presents for the session, subject to any deliberately configured mechanisms such as local-AS behavior. Because the two sides do not agree on R2's presented AS in this configuration, the session remains down until the remote-AS expectation or R2's presented local AS is corrected. Cisco's BGP configuration documentation describes defining BGP neighbors with the `neighbor ... remote-as` command: [Cisco IOS XE BGP Basic Configuration Guide](#). The BGP protocol's OPEN-message AS field is also defined in [RFC 4271, Section 4.2](#).



Refer to the exhibit. A large service provider is migrating device management from Layer 2 VLAN-based to Layer 3 IP-based solution. An engineer must configure the ISIS solution with these requirements:

Network management server IP 172.16.10.1 must be advertised from the core MPLS network to the regional domain.

The external network 202.10.0.0/24 must not establish ISIS peering with the R5 router.

The regional network must prevent sending unnecessary hello packets and flooding the routing tables of the R5 router.

Which two ISIS parameters must be implemented to meet these requirements? (Choose two.)

- A. LSP lifetime maximum
- B. advertise-passive-only
- C. overload bit passive
- D. attached bit on ISIS instance
- E. passive-interface Loopback0

ANSWER: B E

Explanation:

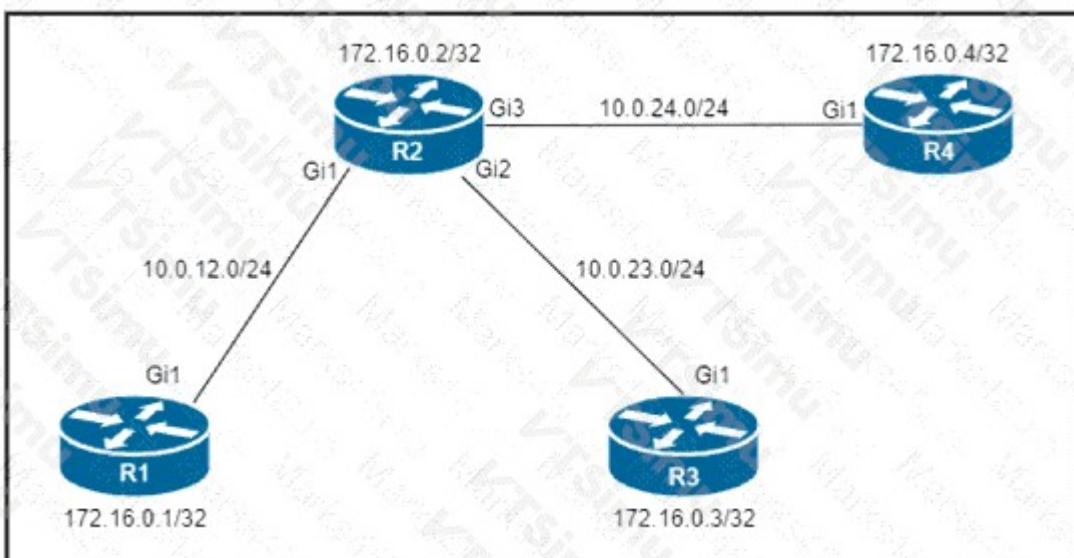
`advertise-passive-only` and `passive-interface Loopback0` together provide the intended IS-IS control-plane design. A passive IS-IS interface does not send IS-IS hello packets and therefore cannot form an IS-IS adjacency on that interface. It can still advertise its associated connected prefix in IS-IS, which allows the management-server address 172.16.10.1 to remain reachable through the IS-IS domain while preventing unwanted peering behavior on the relevant management or external-facing connection.

The `advertise-passive-only` setting restricts IP prefix advertisement to prefixes associated with passive interfaces. This is useful in a service-provider design where only selected infrastructure or management prefixes, such as loopback addresses, need to be distributed. Limiting advertisements in this way avoids injecting unnecessary transit-link prefixes toward R5, reducing the IP routing information that R5 must install and process. The resulting design advertises the required management reachability while minimizing hello traffic and unnecessary route distribution at the regional-domain boundary.

Cisco documents passive-interface behavior for IS-IS and the associated suppression of adjacency formation in its [IS-IS configuration guide](#). Cisco IOS XR command details for IS-IS address-family controls are available in the [Cisco IOS XR IS-IS implementation guide](#).

QUESTION NO: 36

Refer to the exhibit.



Which configuration must be applied to each of the four routers on the network to reduce LDP LIB size and advertise label bindings for the /32 loopback IP space only?

- ☐ config t
ip prefix-list LOOPBACKS seq 5 permit 0.0.0.0/0 le 32
mpls ldp label
allocate global prefix-list LOOPBACKS
end
- ☐ config t
access-list 10 permit 172.16.0.0 0.0.0.7
access-list 20 permit 10.0.0.0 0.0.31.255
no mpls ldp advertise-labels
mpls ldp advertise-labels for 10 to 20
end
- ☐ config t
access-list 10 permit 172.16.0.0 0.0.0.7
access-list 20 permit 172.16.0.0 0.0.0.7
no mpls ldp advertise-labels
mpls ldp advertise-labels for 10 to 20
end
- ☐ config t
mpls ldp label
allocate global host-routes
end

- A. Option A
- B. Option B
- C. Option C
- D. Option D
- E. ip prefix-list LOOPBACKS seq 5 permit 0.0.0.0/0 ge 32 le 32
mpls ldp label
allocate global prefix-list LOOPBACKS

ANSWER: E

Explanation:

The configuration using a prefix list that matches only host routes, followed by LDP global label-allocation filtering, is correct. `ip prefix-list LOOPBACKS seq 5 permit 0.0.0.0/0 ge 32 le 32` matches IPv4 prefixes whose length is exactly 32 bits. Therefore, it selects loopback host routes while excluding all less-specific infrastructure, link, and aggregate routes. Applying that list through `mpls ldp label allocate global prefix-list LOOPBACKS` changes LDP's label-allocation policy so the router creates and advertises label bindings only for FECs permitted by the prefix list. Applying the same policy consistently on every router limits the set of label bindings exchanged throughout the MPLS domain, reducing each router's Label Information Base consumption while preserving MPLS forwarding to the loopback addresses commonly used as router IDs, LDP transport addresses, and BGP next hops. The `ge 32 le 32` qualifiers are essential because they constrain the match to precisely /32 routes rather than permitting every IPv4 prefix under 0.0.0.0/0. Cisco documents prefix-list prefix-length matching and LDP label-allocation filtering in its [IP Routing command reference](#) and [MPLS LDP configuration guide](#).

QUESTION NO: 37

```
PE-A#config t
PE-A(config)#interface FastEthernet0/0
PE-A(config-if)#ip ospf message-digest-key 1 md5 44578611
PE-A(config-if)#ip ospf authentication message-digest

PE-B#config t
PE-B(config)#interface FastEthernet0/0
```

Refer to the exhibit. An engineer wants to authenticate the OSPF neighbor between PE-A and PE-B using MD5.

Which command on PE-B successfully completes the configuration?

- A. PE-B(config-if)#ip ospf message-digest-key 1 md5 44578611 PE-B(config-if)#ip ospf authentication null
- B. PE-B(config-if)#ip ospf message-digest-key 1 md5 44578611
PE-B(config-if)#ip ospf authentication key-chain 44578611
- C. PE-B(config-if)#ip ospf message-digest-key 1 md5 44568611 PE-B(config-if)#ip ospf authentication null
- D. PE-B(config-if)#ip ospf message-digest-key 1 md5 44578611 PE-B(config-if)#ip ospf authentication message-digest

ANSWER: D

Explanation:

PE-B(config-if)#ip ospf message-digest-key 1 md5 44578611 PE-B(config-if)#ip ospf authentication message-digest correctly enables OSPFv2 cryptographic (MD5) authentication on the PE-B interface and defines the MD5 credential used in OSPF hello packets and other OSPF protocol packets. The `ip ospf message-digest-key` command configures key ID 1 with the clear-text key value 44578611; IOS derives and includes the MD5 digest when sending authenticated OSPF packets. The `ip ospf authentication message-digest` command is the interface-level setting that instructs OSPF to use message-digest authentication rather than simple password authentication or no authentication.

For the adjacency to form, the connected PE-A interface must use compatible OSPF authentication settings: MD5 authentication must be enabled, and the peer must have a matching key ID and secret for the active key. Once both sides have matching settings, received OSPF packets can be validated successfully and normal neighbor establishment can proceed. Cisco documents interface-based OSPF authentication and the message-digest key configuration in its [OSPF Authentication Configuration Guide](#) and the [Cisco IOS IP Routing: OSPF Command Reference](#).

QUESTION NO: 38

Which two PHY modes are available to implement an IOS XR Gigabit Ethernet interface interface? (Choose two.)

- A. SONET
- B. MAN
- C. WDM
- D. LAN
- E. WAN

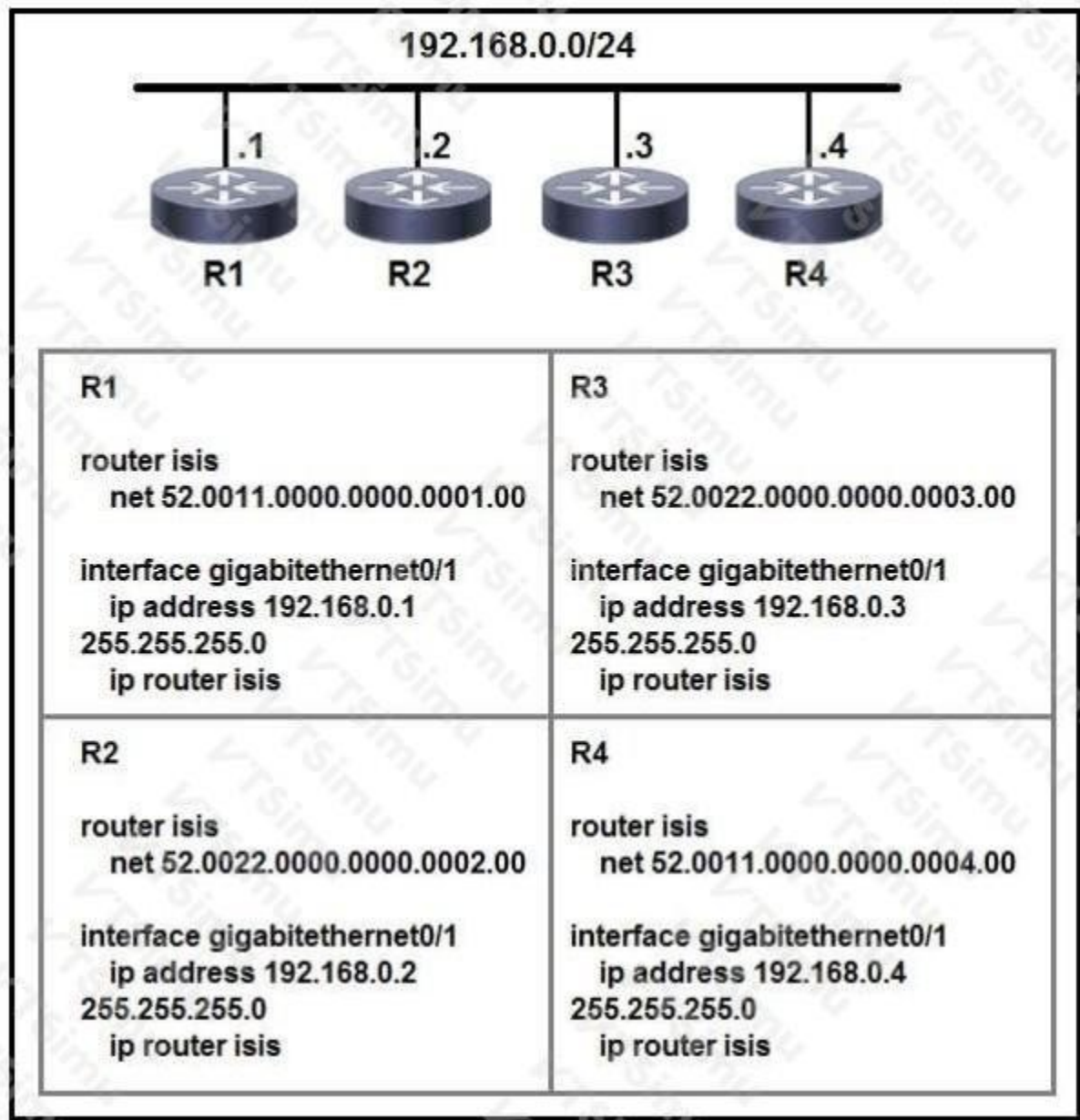
ANSWER: D E

Explanation:

LAN and **WAN** are the applicable Ethernet PHY operating modes. LAN PHY is the native Ethernet physical-layer mode used for Ethernet transport at the line rate defined for the interface. WAN PHY is the Ethernet mode designed for operation in

wide-area transport environments, using a rate and framing approach intended to interoperate with SONET/SDH-oriented transport infrastructure. Cisco IOS XR platform interface documentation and controller command references distinguish LAN and WAN PHY operation where the installed Ethernet hardware supports selectable PHY modes. This terminology also aligns with the IEEE 10 Gigabit Ethernet architecture, which specifies LAN and WAN physical-layer variants: LAN PHY is used for conventional Ethernet applications, while WAN PHY supports carriage over WAN transport systems. The selected PHY mode must be supported by the specific line card or port and is configured in the IOS XR interface/controller context appropriate to that hardware. See Cisco's [ASR 9000 installation and configuration guides](#) and the IEEE [802.3ae 10 Gigabit Ethernet project information](#) for the LAN PHY and WAN PHY model.

QUESTION NO: 39



Refer to the exhibit. Which two statements about the IS-IS topology are true? (Choose two.)

- A. R1 and R4 are Level 2 neighbors.
- B. All four routers are operating as Level 1-2 routers.
- C. All four routers are operating as Level 2 routers only.
- D. All four routers are operating as Level 1 routers only.
- E. R1 and R2 are Level 2 neighbors.

ANSWER: A B**Explanation:**

R1 and R4 are Level 2 neighbors because the topology shows an adjacency between these routers and both ends support Level 2 operation. An IS-IS Level 2 adjacency is formed by routers that share a common data link and have Level 2 enabled; unlike Level 1 adjacency formation, it is not restricted by an IS-IS area match. This permits Level 2 to provide the interarea backbone connectivity in an IS-IS domain.

All four routers are operating as Level 1-2 routers. A Level 1-2 router participates in both IS-IS levels: it maintains Level 1 knowledge for its local area and Level 2 knowledge for backbone and interarea reachability. Consequently, these routers can establish the appropriate Level 1 relationships within an area while also participating in the Level 2 topology. Cisco documents that a router configured for both levels functions as a Level 1-2 router and can maintain separate Level 1 and Level 2 link-state databases. This dual-level role is the topology model represented in the exhibit. See Cisco's [IS-IS overview documentation](#) and [Cisco's IS-IS technical overview](#).

QUESTION NO: 40

What are two factors to consider when implementing NSR High Availability on an MPLS PE router? (Choose two.)

- A. It consumes more memory and CPU resources than NSF
- B. It operates normally without NSR support on the PE peers.
- C. It requires all PE-CE sessions to support NSR
- D. It requires routing protocol extensions
- E. It cannot sync state information across redundant RPs

ANSWER: A B**Explanation:**

"It consumes more memory and CPU resources than NSF" is correct because Nonstop Routing maintains synchronized control-plane and routing-protocol state between the active and standby route processors. This state replication allows the standby processor to take over routing operations with minimal disruption following a route-processor switchover. Maintaining protocol state, RIB-related information, and other control-plane data necessarily imposes additional memory and processing overhead compared with approaches that focus primarily on preserving forwarding-plane operation. Capacity planning for the PE must therefore account for that overhead, particularly when the router carries large BGP, VPN, or IGP tables.

"It operates normally without NSR support on the PE peers." is also correct. NSR is a local high-availability capability of the redundant route processors in the PE router; its purpose is to synchronize state internally so that a switchover is largely transparent to adjacent routers. Neighboring PEs and CEs do not need to implement NSR themselves for the local router's NSR function to work. Cisco documents NSR as stateful switchover support for routing protocols and distinguishes it from mechanisms that depend on neighbor restart awareness. See Cisco's [IOS XR Nonstop Routing configuration documentation](#) and its [system redundancy guide](#).

QUESTION NO: 41

What is the function of the FEC field within the OTN signal structure?

- A. It allows the sending devices to apply QoS within the OTN forwarding structure.
- B. It allows source nodes to discard payload errors before transmitting data on the network.
- C. It allows receivers to correct errors upon data arrival.
- D. It allows deep inspection of data payload fields.

ANSWER: C

Explanation:

It allows receivers to correct errors upon data arrival. Forward error correction (FEC) is a mechanism in the Optical Transport Network (OTN) frame structure that adds redundant coding information to the transmitted signal. At the receiving end, this information enables the receiver to detect and correct a defined number of bit errors introduced by optical impairments, such as attenuation, dispersion, and noise, without requiring retransmission from the source. This capability improves the error performance and reach of high-speed optical transport links while preserving efficient use of the transport path. In the ITU-T OTN framing architecture, FEC is associated with the optical transport unit (OTUk) and is processed at OTN termination points: the transmitting endpoint generates the FEC data and the receiving endpoint uses it to perform error correction. FEC performance can also be monitored through corrected and uncorrectable error statistics, which helps operators assess optical-link quality and determine whether maintenance is needed. Cisco's OTN documentation describes FEC as an OTU function used to improve bit-error-rate performance, while the ITU-T G.709 recommendation defines the OTN interfaces and frame structure that carry this information. See [ITU-T Recommendation G.709](#) and [Cisco OTN Configuration Guide](#).

QUESTION NO: 42

Which two statements about LDP targeted sessions are true? (Choose two.)

- A. It can establish an LDP session between non-directly connected LSRs.
- B. It uses targeted Hellos addressed to a specified peer.
- C. It requires both peers to share the same Layer 2 segment.
- D. It distributes labels only for VPNv4 routes.
- E. It replaces the TCP connection used by an LDP session.

ANSWER: A B

Explanation:

An LDP targeted session is established between LSRs that are not directly connected. Unlike a link LDP session, which is discovered with link Hellos sent to the all-routers multicast address on an MPLS-enabled interface, a targeted session uses targeted Hellos addressed to a specific peer. This allows LDP session establishment across an IP-routed network.

Targeted LDP is commonly used for applications such as LDP-over-RSVP-TE, interarea LDP transport, and certain Layer 3 VPN designs. A targeted Hello is sent to UDP port 646 and identifies the intended remote LSR. The remote router must accept targeted Hellos and establish the TCP LDP session using the transport addresses negotiated during discovery. The LDP targeted-discovery procedures are defined in [RFC 5036](#).

QUESTION NO: 43

An engineer implemented LDP protocol on the ISP network. The engineer must ensure that there are no packet loss issues when IGP and LDP protocols are not synchronized. Which configuring must the engineer implement so that the IGP routing protocol will wait until LDP convergence is completed?

- A. Disable IP CEF routers running LDP and enable LDP protocol.
- B. Configure MPLS LDP IGP synchronization on the network.
- C. Configure LDP sessions protection on the network.
- D. Disable MPLS LDP IGP synchronization on the network.

ANSWER: B

Explanation:

Configure MPLS LDP IGP synchronization on the network. This feature prevents an IGP from selecting an MPLS-capable link for transit traffic before LDP has established label exchange with the neighboring router on that link. Without synchronization, the IGP can advertise the link as the preferred path as soon as the IGP adjacency is up, while LDP may still be converging. Traffic forwarded toward that next hop can then arrive without an appropriate label forwarding entry, creating a transient packet-loss condition.

With LDP-IGP synchronization enabled, the router advertises a maximum IGP metric for the interface until the LDP session is operational and labels are available. After LDP synchronization completes, the interface returns to its normal configured IGP metric and can safely become a preferred path. This behavior is particularly important in service-provider MPLS cores, where the IGP supplies reachability and LDP supplies the transport labels used to forward labeled traffic. Cisco documents this mechanism as LDP-IGP synchronization and supports its use with IGP protocols such as OSPF and IS-IS. See Cisco's [MPLS LDP IGP Synchronization documentation](#) and the [Cisco IOS MPLS LDP Configuration Guide](#).

QUESTION NO: 44

What are two purposes of a data-modeling language such as YANG? (Choose two.)

- A. defining a hierarchical data scheme
- B. modeling notifications, data configuration, and data status
- C. defining a relational-model data scheme
- D. defining a semi-structured data scheme
- E. defining API calls, JSON format, and XML structure

ANSWER: A B

Explanation:

YANG is a data-modeling language designed to model the configuration and operational state of network devices and services in a structured, implementation-independent way. A central purpose is **defining a hierarchical data scheme**: YANG organizes modeled information as a tree of nodes, such as containers, lists, leafs, and leaf-lists. This hierarchy represents how related network configuration and operational data is logically arranged, allowing clients and servers to share a consistent schema.

YANG also supports **modeling notifications, data configuration, and data status**. A YANG module can define configuration data that a management client may set, state data that reports operational or read-only information, and notifications that communicate asynchronous events from a server. These capabilities let the same model describe both the desired configuration and the observable behavior of a network service or device. Protocols such as NETCONF and RESTCONF use YANG models to validate, expose, retrieve, and manipulate this modeled data. RFC 7950 defines YANG's hierarchical data model and its constructs for configuration, state data, and notifications. See [RFC 7950: The YANG 1.1 Data Modeling Language](#) and Cisco's [Introduction to YANG](#).

QUESTION NO: 45

What is a key capability of EVPN?

- A. providing built-in support for native integration with third-party cloud platforms
- B. supporting multitenancy for Layer 2 and Layer 3 VPN services
- C. establishing IPv6 overlay networks over IPv4 networks
- D. enabling real-time voice and video traffic optimization

ANSWER: B

Explanation:

supporting multitenancy for Layer 2 and Layer 3 VPN services is a key EVPN capability. Ethernet VPN (EVPN) is a standards-based control plane, typically using BGP, for distributing MAC and IP reachability information across an MPLS or VXLAN-based overlay. It enables separate customer or tenant networks to share a common provider or data-center fabric while maintaining logical isolation through constructs such as EVPN instances, VLANs or bridge domains, VRFs, and route targets.

For Layer 2 services, EVPN provides multipoint Ethernet connectivity and advertises endpoint MAC addresses through BGP rather than relying solely on data-plane flooding. For Layer 3 services, EVPN can integrate IP VPN/VRF routing and advertise IP-prefix reachability, enabling tenant-specific routed connectivity in the same EVPN architecture. This unified L2 and L3 multitenant model is especially useful in service-provider and data-center networks because it provides scalable segmentation, controlled route distribution, and efficient endpoint mobility handling. Cisco describes EVPN as a BGP-based control plane that delivers Layer 2 and Layer 3 VPN services over an overlay infrastructure. See Cisco's [VXLAN BGP EVPN configuration guide](#) and the IETF's [RFC 7432](#) for the EVPN control-plane model.

QUESTION NO: 46

A network architect at ISP must implement a loop-avoidance mechanism in the organization's ring network topology. The architect decided to implement ITU-T G.8032. Intermittent link flaps within two seconds should be ignored. Which two configurations must the architect apply? (Choose two.)

- A. ring-protection g8032 profile vlan port-channel none
- B. link-protection group management vlan 1
- C. timer hold-off 2
- D. continuity-check direction down 2
- E. ethernet ring g8032 profile ethernet

ANSWER: C E

Explanation:

`ethernet ring g8032 profile ethernet` establishes the Ethernet Ring Protection Switching configuration context for a G.8032 Ethernet ring. G.8032 provides loop prevention and sub-50-ms protection switching for Ethernet ring topologies by maintaining a controlled blocked ring link during normal operation and changing ring-port states when a valid fault is detected. A profile is the container in which the ring's protection behavior, ports, and timers are defined.

`timer hold-off 2` configures a two-second hold-off interval. The hold-off timer delays fault processing after a link defect is first detected. Consequently, a transient defect or link flap that clears during that two-second interval does not trigger an Ethernet Ring Protection Switching event. This is specifically useful where brief physical-layer instability should not cause unnecessary protection switching, ring-port state changes, or control-message activity. The hold-off setting is applied within the G.8032 profile so it governs fault handling for that ring's configured protection domain. ITU-T defines the Ethernet ring protection mechanism and its operational model in [ITU-T Recommendation G.8032](#); Cisco's Ethernet ring protection configuration guidance is available in the [Cisco IOS XE Ethernet Ring Protection documentation](#).

QUESTION NO: 47

Which two tasks must an engineer perform when implementing LDP NSF on the network? (Choose two.)

- A. Disable Cisco Express Forwarding.
- B. Enable NSF for EIGRP.
- C. Enable NSF for the link-state routing protocol that is in use on the network.
- D. Implement direct connections for LDP peers.
- E. Enable NSF for BGP.

ANSWER: C D

Explanation:

LDP Nonstop Forwarding is intended to preserve MPLS forwarding through a route-processor switchover by allowing forwarding information to remain usable while the control plane recovers. The required implementation task is to enable NSF for the link-state routing protocol deployed in the network, such as OSPF or IS-IS. LDP relies on the IGP's NSF behavior to preserve reachability and to allow the restarted router to recover its routing state without unnecessarily disrupting label-switched forwarding. Cisco documents LDP NSF as depending on NSF support in the underlying IGP.

The LDP peers must also be directly connected. This topology constraint ensures that the LDP session and its associated label bindings can be recovered predictably during the NSF event. Directly connected LDP neighbors use the link-local peer relationship needed by the feature's recovery model. With these two conditions in place, existing MPLS forwarding can continue while the control plane and LDP state are restored after a switchover. Cisco's MPLS LDP configuration guidance is available in the [Cisco IOS XE MPLS LDP Configuration Guide](#); background on the LDP graceful-restart protocol is defined in [RFC 3478](#).

QUESTION NO: 48

After recent network outages and customer complaints, a large Tier 1 service provider has asked a network engineer to implement MPLS OAM policy on all Provider Edge devices in the MPLS network. To improve network performance and reliability, the policy must be able to detect 2% packet drops and a delay of 0.05 seconds. The solution must be based on RFC 6374 and RFC 4379 standards for MPLS OAM. The CCM to detect connectivity issues has already been configured with interval 1000. Which two actions must the engineer take to achieve the goal? (Choose two.)

- A. Set loopback delay-frequency 50 and drop-threshold 0.2 on the PE routers.
- B. Implement mpls ldp tracing on the PE and P routers.
- C. Implement mpls oam policy oam on the PE routers.
- D. Implement mpls ip oam map-policy on the PE routers.
- E. Set loss-threshold 2 and delay-threshold 50 on the P routers.

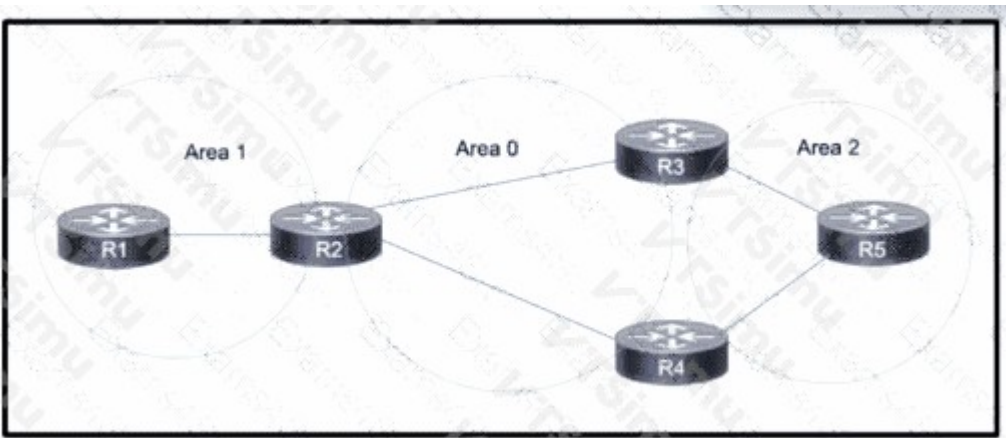
ANSWER: A C

Explanation:

Set loopback delay-frequency 50 and drop-threshold 0.2 on the PE routers. configures the operational measurement parameters needed by the MPLS OAM policy. The loopback delay frequency supports active delay testing at the required granularity, while the drop threshold establishes the loss trigger used to identify the specified packet-drop condition. These measurements complement the already configured 1000-ms continuity-check interval: continuity monitoring detects connectivity defects, whereas delay and loss measurements provide the performance information required to identify service degradation.

Implement mpls oam policy oam on the PE routers. applies the MPLS OAM policy at the Provider Edge, where service LSPs are originated or terminated and where the provider can measure the customer-facing transport service end to end. RFC 4379 defines MPLS LSP ping and traceroute mechanisms used for MPLS fault isolation and path verification. RFC 6374 defines MPLS-TP OAM loss and delay measurement functions, including the framework for performance monitoring. Applying the policy on each PE therefore enables standards-based monitoring of the relevant MPLS services and allows the configured thresholds to generate actionable performance detection. See [RFC 4379](#) and [RFC 6374](#).

QUESTION NO: 49



Refer to the exhibit. OSPF is running in the given network. R1 was previously running EIGRP on a standalone network, and a network engineer just migrated it to OSP

F.

The R1 peering with R2 has come up.

Which additional action must the engineer take so that R2 receives routes from Area 1 and Area 0 but is prevented from receiving routes from Area 2?

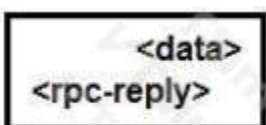
- A.** Implement a route map that matches R2's interface to R3 and R4, and then attach the route map to an inbound distribute list on R2.
- B.** Implement a route map that matches the routes in an ACL from Area 2, and then attach the route map to an inbound distribute list on R2.
- C.** Implement a route map that matches the routes from Area 2 in a prefix list, and then attach the route map to an outbound distribute list on R2.
- D.** Implement a route map that matches R1's next-hop address, and then attach the route map to an outbound distribute list on R2.

ANSWER: B

Explanation:

Implementing a route map that matches the routes in an ACL from Area 2, and then attaching the route map to an inbound distribute list on R2 provides the required route-selection behavior. The ACL identifies the specific Area 2 destination prefixes that must not be installed by R2. The route map references that ACL and can deny those matching prefixes while permitting all remaining OSPF routes. Applied inbound under the OSPF process on R2, the distribute list filters OSPF routes as R2 considers them for installation in its IP routing table. Consequently, Area 1 and Area 0 prefixes remain eligible for installation, while the selected Area 2 prefixes are withheld from R2's routing table. This is a local route-filtering mechanism: OSPF link-state advertisements are still exchanged and retained as required for OSPF's link-state operation, but the policy controls which computed routes R2 uses for forwarding. Cisco documents OSPF distribute-list support for inbound filtering and route-map use in its [OSPF Route Filtering Configuration Guide](#). The underlying OSPF behavior of distributing topology information and independently calculating routes is defined in [RFC 2328](#).

QUESTION NO: 50



Refer to the exhibit. This output is included at the end of an output that was provided by a device using NETCONF.

What does the code show?

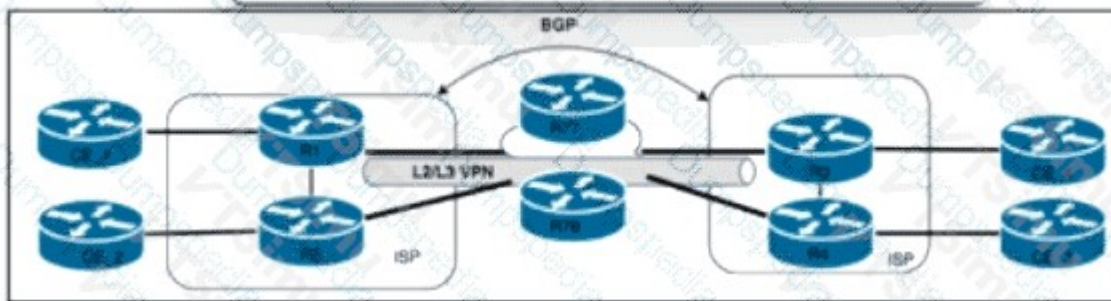
- A. It shows that the full configuration is being modeled by YANG.
- B. It shows NETCONF uses remote procedure calls.
- C. It shows the hostname of the device as rpc-reply.
- D. It shows that the running configuration is blank.

ANSWER: B

Explanation:

It shows NETCONF uses remote procedure calls. NETCONF exchanges are structured as RPC transactions: a client sends an `<rpc>` request containing an operation, and the NETCONF server returns an `<rpc-reply>` message. The ending XML shown in this type of device output identifies the completion of that reply message and therefore reflects NETCONF's RPC-based protocol model. The reply is correlated to the request through the request's `message-id`, allowing a client to associate a returned result, configuration data, or status with the operation it invoked. This RPC model is fundamental to NETCONF operations such as retrieving configuration with `<get-config>`, retrieving state and configuration data with `<get>`, or modifying configuration using edit operations. An `<rpc-reply>` wrapper is protocol messaging, not a hostname or a direct statement about whether the running configuration contains data. The NETCONF base specification defines `<rpc>` and `<rpc-reply>` as the request and response elements used for protocol operations. See [RFC 6241, Network Configuration Protocol \(NETCONF\)](#) and Cisco's [NETCONF documentation](#).

QUESTION NO: 51



Refer to the exhibit. Company A's network architect is implementing the capabilities that are described in RFC 4379 to provide new options for troubleshooting the company's MPLS network. MPLS LSPs are signaled via the BGP routing protocol with IPv4 unicast prefix FECs. Due to company security policies, access is controlled on the control plane and management plane. Which two tasks must the architect perform to achieve the goal? (Choose two.)

- A. Open UDP port 3503.
- B. Open MCP port 3505.
- C. Implement OAM on all MPLS-enabled routers.
- D. Implement LSP Ping on PE routers only.
- E. Implement VCCV on core routers only.

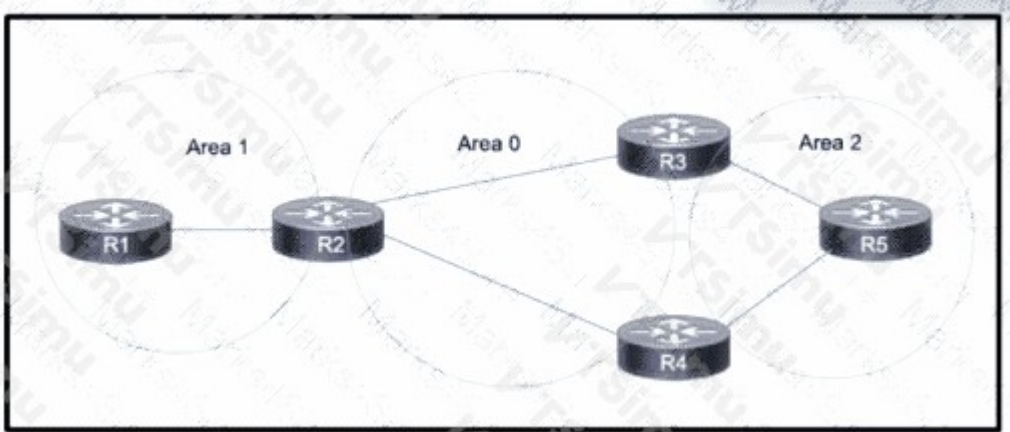
ANSWER: A C

Explanation:

RFC 4379 defines MPLS LSP Ping and MPLS traceroute, which provide data-plane verification and path-discovery functions for MPLS label-switched paths. **Open UDP port 3503.** is required because an MPLS Echo Request is encapsulated in UDP and uses destination port 3503. Control-plane protection and management-plane filtering therefore must permit this traffic where LSP Ping or traceroute requests and responses are processed. The use of BGP-signaled IPv4 unicast prefix FECs identifies the LSP being tested, but it does not change the RFC 4379 UDP transport-port requirement.

Implement OAM on all MPLS-enabled routers. is also required for complete RFC 4379 troubleshooting behavior across the MPLS domain. MPLS traceroute relies on LSRs along the label-switched path recognizing an expired MPLS TTL and generating the appropriate MPLS Echo Reply. Enabling the relevant MPLS OAM functionality throughout the participating MPLS forwarding path allows the architect to validate LSP continuity and identify each responding hop, rather than limiting diagnostics to only an ingress or egress device. RFC 4379 specifies the Echo Request UDP encapsulation and traceroute processing model; Cisco documents LSP Ping as an MPLS OAM mechanism for validating forwarding paths. See [RFC 4379](#) and [Cisco MPLS OAM Configuration Guide](#).

QUESTION NO: 52



Refer to the exhibit. OSPF is running in the given network. R1 was previously running EIGRP on a standalone network, and a network engineer just migrated it to OSPF. The R1 peering with R2 has come up.

Which additional action must the engineer take so that R2 receives routes from Area 1 and Area 0 but is prevented from receiving routes from Area 2?

- A. Implement a route map that matches R2 ' s interface to R3 and R4, and then attach the route map to an inbound distribute list on R2.
- B. Implement a route map that matches the routes in an ACL from Area 2, and then attach the route map to an inbound distribute list on R2.
- C. Implement a route map that matches the routes from Area 2 in a prefix list, and then attach the route map to an outbound distribute list on R2.
- D. Implement a route map that matches R1 ' s next-hop address, and then attach the route map to an outbound distribute list on R2.

ANSWER: B

Explanation:

Implement a route map that matches the routes in an ACL from Area 2, and then attach the route map to an inbound distribute list on R2. An inbound OSPF distribute list controls whether routes learned through OSPF are installed in the local routing table. It does not prevent OSPF adjacencies from forming or stop the underlying link-state advertisements from being flooded; instead, it provides the required route-installation policy at R2. The ACL identifies the prefixes that belong to Area 2, and the route map uses that match to deny installation of those prefixes while permitting the Area 1 and Area 0 prefixes that R2 must retain. The distribute list must be applied inbound because the policy requirement concerns routes R2 learns and places into its own routing table after receiving OSPF information from its neighbor. This approach is particularly appropriate when the goal is selective route filtering on one router rather than changing OSPF's area-wide interarea route advertisement behavior. Cisco documents that OSPF distribute lists can filter routes from being installed in the routing table, and the base OSPF specification distinguishes the link-state database from the route-calculation and route-installation process. See the [Cisco IOS XE OSPF Configuration Guide](#) and [RFC 2328](#).

QUESTION NO: 53

A company needs to improve the use of the network resources that is used to deploy internet access service to customers on separate backbone and internet access network. Which two major design models should be used to configure MPLS L3VPNs and internet service in the same MPLS backbone? (Choose two.)

- A. Carriage of full internet routes in a VPN, in the case of internet access VPNS
- B. Internet routing through global routing on a PE router.
- C. Internet access routing as another VPN in the ISP network.
- D. Internet access through leaking of internet routed from the global table into the L3VPN VRF
- E. Internet access for global routing via a separate interface in a VRF

ANSWER: B C

Explanation:

The major MPLS Layer 3 VPN Internet-access models are **Internet routing through global routing on a PE router.** and **Internet access routing as another VPN in the ISP network.** In the global-routing model, the PE uses its global routing table for Internet reachability while maintaining customer routes in their respective VRFs. Customer Internet-bound traffic is forwarded from the VPN context toward the PE's Internet-facing global-routing infrastructure, allowing the provider to use a shared Internet connection and common Internet routing resources.

In the Internet-VPN model, the provider creates a dedicated VRF/VPN that represents the Internet service. Customer VPNs import an Internet default route, selected Internet prefixes, or both from that Internet VPN, and Internet return traffic is directed through the same VPN-based architecture. This approach keeps Internet routing logically separated from the PE global table and uses standard MPLS VPN route-target import and export policy to control which customer VPNs receive Internet access.

Both models permit Internet and customer L3VPN services to share the provider MPLS backbone while retaining controlled routing separation. Cisco's MPLS VPN architecture is based on VRF-separated routing and controlled route distribution, as described in [RFC 4364](#) and Cisco's [MPLS VPN configuration guidance](#).

QUESTION NO: 54

A network engineer is testing an automation platform that interacts with Cisco networking devices via NETCONF over SSH. In accordance with internal security requirements:

NETCONF sessions are permitted only from trusted sources in the 172.16.20.0/24 subnet.

CLI SSH access is permitted from any source.

Which configuration must the engineer apply on R1?

- A. hostname R1
- B. hostname R1
- C. hostname R1
- D. hostname R1
- E. `netconf-yang ip access-list extended NETCONF_SSH_TRUSTED permit tcp 172.16.20.0 0.0.0.255 any eq 830 deny tcp any any eq 830 permit ip any any interface ip access-group NETCONF_SSH_TRUSTED in`

ANSWER: E

Explanation:

The configuration that enables NETCONF/YANG and applies an inbound extended ACL to every ingress interface is correct because NETCONF over SSH normally uses TCP port 830, while ordinary interactive SSH access uses its separate SSH service port. The ACL explicitly permits TCP/830 only when the source is within 172.16.20.0/24, denies TCP/830 from every

other source, and then permits all remaining IP traffic. Consequently, SSH traffic used for CLI administration remains unrestricted by source, satisfying the requirement that CLI SSH access be allowed from anywhere. Applying the ACL inbound on each interface through which management traffic can arrive ensures that unauthorized NETCONF connection attempts are discarded before they reach the NETCONF service. The `netconf-yang` global command enables the IOS XE NETCONF/YANG agent; by default, NETCONF over SSH uses the IANA-assigned NETCONF SSH port, TCP/830. The port assignment and SSH transport behavior are defined in [RFC 6242](#). Cisco's IOS XE NETCONF/YANG configuration guidance is available in the [Cisco IOS XE Programmability Configuration Guide](#).

QUESTION NO: 55

Which two actions are performed by Reverse Path Forwarding checks in multicast routing? (Choose two.)

- A. It determines the expected incoming interface by looking up the multicast source in the unicast routing table.
- B. It drops a multicast packet that arrives on a non-RPF interface.
- C. It selects the designated router on a shared Ethernet segment.
- D. It assigns the rendezvous point for a multicast group.
- E. It converts a shared tree into a source tree.

ANSWER: A B

Explanation:

Multicast Reverse Path Forwarding uses the unicast routing table to determine the expected upstream interface toward a multicast source. When a multicast packet arrives, the router compares the receiving interface with the interface that the unicast routing table identifies as the best path back to the source. If the packet arrives on that expected interface, it passes the RPF check and is eligible for multicast forwarding.

If the packet arrives on an interface other than the RPF interface, the router normally drops it. This prevents duplicated multicast packets from being forwarded around loops. RPF behavior is a fundamental multicast forwarding mechanism and is described in Cisco multicast routing documentation, including the [Cisco multicast RPF troubleshooting guide](#).

QUESTION NO: 56

```
R1
interface fastethernet1/0
  ip address 192.168.2.14 255.255.255.0
  ip ospf message-digest-key 1 md5 cisco
  ip ospf authentication message-digest
```

Refer to the exhibit. Which condition must be met by the OSPF peer of router R1 before the two devices can establish communication?

- A. The OSPF peer must use clear-text authentication.
- B. The OSPF peer must be configured as an OSPF stub router.
- C. The interface on the OSPF peer may have a different key ID, but it must use the same key value as the configured interface.
- D. The interface on the OSPF peer must use the same key ID and key value as the configured interface.

ANSWER: D

Explanation:

The interface on the OSPF peer must use the same key ID and key value as the configured interface. OSPF neighbors exchange Hello packets before they can form an adjacency, and authentication is verified as part of processing those packets. With cryptographic (MD5) authentication, the sending router includes a key ID in the OSPF authentication data and calculates an MD5 digest using the configured secret key. The receiving router uses that key ID to select its locally configured key and validates the digest. Therefore, the peer must have the corresponding key ID configured with the identical key value for authentication to succeed and for the routers to communicate as OSPF neighbors. Key IDs also support orderly key rollover: multiple keys can be configured, but any key actively used in received packets must be present locally under the matching identifier and secret. This behavior is defined in the OSPFv2 cryptographic authentication procedure in [RFC 2328, Appendix D](#). Cisco's OSPF documentation also describes configuring message-digest keys on participating interfaces for OSPF authentication: [Cisco IOS OSPFv2 Authentication Configuration Guide](#).

QUESTION NO: 57

How does Inter-AS Option-A function when two PE routers in different autonomous systems are directly connected?

- A. The two routers share all Inter-AS VPNv4 routes and redistribute routes within an IBGP session to provide end-to-end reach.
- B. The two routers establish an MP-EBGP session to share their customers' respective VPNv4 routes.
- C. The two routers treat one another as CE routers and advertise unlabeled IPv4 routes through an EBGP session.
- D. The two routers share VPNv4 routes over a multihop EBGP session and set up an Inter-AS tunnel using one another's label.

ANSWER: C

Explanation:

Inter-AS Option A uses a back-to-back VRF design at the autonomous-system boundary. The directly connected provider-edge routers do not exchange VPNv4 Network Layer Reachability Information across the AS boundary. Instead, each router places the inter-AS-facing connection in a VRF and regards the neighboring provider-edge router as though it were a customer-edge router for that VRF. The routers then exchange ordinary IPv4 customer prefixes, typically using an external BGP peering. Because the inter-AS eBGP exchange carries IPv4 VPN customer reachability rather than VPNv4 routes, no VPN label is advertised across that peering; each provider network applies its own local MPLS/VPN label handling on its side of the boundary. This model provides clear separation of the two provider MPLS cores and is commonly used where the providers do not want to exchange VPNv4 routes or coordinate MPLS label distribution between their networks. Cisco describes this approach as connecting the two ASBR/PE devices through VRFs and using per-VRF routing between them. See Cisco's [Interautonomous System MPLS VPN configuration guide](#) and the [Cisco MPLS Layer 3 VPN documentation](#).

QUESTION NO: 58

Refer to the exhibit. A network engineer is configuring Ethernet Layer 2 service to connect CE2 and CE3 for video and data application sharing with these requirements:

A point-to-point cross-connect service must be established between 10.10.10.10 and 20.20.20.20.

PE1 and PE2 must learn neighbors dynamically in PW 10.

Which configuration must be implemented on PE1 to meet the requirements?

- A. `xconnect group XCON1 p2p XCON1_PE1PE2 interface GigE 0/0 pw-class Path1 backup neighbor 20.20.20.20 pw-id 10`
- B. `12vpn xconnect group XCON1 interface GigE 0/0 interface GigE 0/1`
- C. `12vpn xconnect group XCON1 p2p XCON1_PE1PE2 interface GigE 0/0 neighbor 20.20.20.20 pw-id 10 mpls static label local 699 remote 890`
- D. `12vpn p2p XCON1_PE1PE2 interface GigE 0/1 neighbor 20.20.20.20 pw-id 10`

E. l2vpn xconnect group XCON1 p2p XCON1_PE1PE2 interface GigE 0/0 neighbor 20.20.20.20 pw-id 10

ANSWER: E

Explanation:

The required configuration is the L2VPN point-to-point xconnect definition that binds the local attachment circuit on GigE 0/0 to a pseudowire neighbor at 20.20.20.20 with pseudowire ID 10. Under IOS XR, the `l2vpn` hierarchy contains an `xconnect` group, and the `p2p` section defines a single point-to-point Ethernet service. The local `interface` command identifies the CE-facing attachment circuit, while `neighbor 20.20.20.20 pw-id 10` creates the remote pseudowire endpoint.

Because no static MPLS labels are configured, pseudowire label binding and signaling are learned dynamically through the MPLS LDP pseudowire signaling process. Both PEs use the same PW ID to associate their respective xconnect endpoints, enabling the Layer 2 service between the two customer-facing interfaces. This is the standard IOS XR model for an LDP-signaled Ethernet-over-MPLS point-to-point service. Cisco's IOS XR L2VPN configuration guidance is available in the [Cisco IOS XR L2VPN Configuration Guide](#). The underlying LDP pseudowire signaling behavior is defined by [RFC 4447](#).

QUESTION NO: 59

A network engineer implements OSPF on the network and then notices frequent faults in traffic between two routers, particularly on interface TenGigE 0/0/1. The engineer decides to configure NetFlow to monitor the issue. Which configuration must the engineer apply to the router?

- A. interface tenGigE 0/0/1 ip flow egress snmp-server community public snmp-server enable traps lsnmp-server host 192.168.1.27 version 2c public
- B. interface tenGigE 0/0/1 ip flow ingress ip flow-export version 9 ip flow-cache timeout active 1
- C. ip flow-export destination 192.168.2.1 2055 ip flow-export source loopback0 ip flow-export version 9 ip flow-cache timeout active 1 ip flow-cache timeout inactive 15 snmp-server ifindex persist
- D. ip flow-export destination 192.168.2.1 2055 ip flow-export source loopback0 ip flow-export version 9 ip flow-cache timeout active 1 ip flow-cache timeout inactive 15 snmp-server ifindex persist interface tenGigE 0/0/1 ip flow ingress

ANSWER: D

Explanation:

The configuration containing both the global NetFlow export settings and `ip flow ingress` under interface `tenGigE 0/0/1` is required. NetFlow does not collect traffic from an interface merely because an export destination and export parameters are configured globally. Applying `ip flow ingress` to the affected interface enables ingress flow accounting, allowing the router to create flow-cache entries for packets received on TenGigE 0/0/1. The global commands then define how those collected records are exported: `ip flow-export destination 192.168.2.1 2055` identifies the collector and UDP port, `ip flow-export source loopback0` supplies a stable source address, and NetFlow version 9 exports template-based records suitable for a collector. The active and inactive cache timers cause long-lived and completed flows to be exported promptly, which is useful while investigating intermittent traffic faults. The interface-index persistence setting preserves SNMP interface-index mappings across reloads, helping monitoring systems maintain consistent interface identification. Cisco's NetFlow documentation describes the need to enable flow collection on an interface and separately configure export parameters: [Cisco IOS NetFlow Version 9 Configuration Guide](#). See also [Cisco IOS NetFlow Command Reference](#).

QUESTION NO: 60

Which two features will be used when defining SR-TE explicit path hops if the devices are using IP unnumbered interfaces? (Choose two.)

- A. router ID
- B. labels

- C. node address
- D. next hop address
- E. output interface

ANSWER: A E

Explanation:

router ID and **output interface** are used to identify an explicit path hop across an IP unnumbered link. Because an unnumbered interface has no unique IPv4 address that can be used as the hop endpoint, the path definition identifies the router owning the link through its stable router ID and identifies the particular egress link through its output-interface identifier. Together, these values uniquely describe the intended adjacency, allowing the SR-TE path computation or signaling process to select the correct link even where normal next-hop IPv4 addressing is unavailable.

This is consistent with the unnumbered-interface representation defined for MPLS traffic-engineering explicit routes: an unnumbered link is represented using the router ID of the router that owns the interface and an interface ID identifying that interface. In operational SR-TE deployments, using a stable router ID also ensures that the node remains consistently identifiable as interface addressing and physical links are managed. Cisco SR documentation describes SR-TE policies as imposing an explicitly selected sequence of segments, while the unnumbered-link encoding supplies the information needed to identify a specific adjacency when an address cannot do so.

References: [RFC 3477: Signalling Unnumbered Links in RSVP-TE](#); [Cisco IOS XR Segment Routing Configuration Guide](#).

QUESTION NO: 61

After you analyze your network environment, you decide to implement a full separation model for Internet access and MPLS L3VPN services.

For which reason do you make this decision?

- A. It enables EGP and IGP to operate independently.
- B. It enables you to choose whether to separate or centralize each individual service.
- C. It is easier to manage a system in which services are mixed.
- D. It requires only one edge router.

ANSWER: A

Explanation:

It enables EGP and IGP to operate independently. A full separation model uses distinct network functions and routing domains for Internet access and MPLS Layer 3 VPN services rather than combining them on the same edge infrastructure. This separation lets the Internet-facing environment run its external-routing policy independently, typically using BGP to exchange routes with upstream providers, peers, or customers. At the same time, the MPLS VPN provider core can maintain its own internal routing design, including its IGP and MPLS transport behavior, without being coupled to Internet-route processing or Internet policy changes.

This operational independence is valuable when the services have different scalability, security, resiliency, policy, or change-control requirements. For example, an Internet-routing event or a modification to external route policy can be contained within the Internet service domain, while the MPLS L3VPN control plane and its customer VPN route distribution remain governed by the provider's VPN architecture. Cisco describes MPLS L3VPNs as a provider-managed service that uses separate VPN routing and forwarding information at provider-edge devices; maintaining clear service boundaries supports this model. See [Cisco's MPLS overview](#) and [Cisco IOS XE MPLS L3VPN configuration guide](#).

QUESTION NO: 62

Which three OSPF parameters must match before two devices can establish an OSPF adjacency? (Choose three.)

- A. IP address
- B. subnet mask
- C. interface cost
- D. process ID
- E. area number
- F. hello timer setting

ANSWER: B E F

Explanation:

For two OSPF routers to establish an adjacency on a common multiaccess network, they must agree on fundamental parameters carried and validated in OSPF Hello packets. **subnet mask** must match because it verifies that both interfaces share the same network segment. This check prevents routers that are configured for inconsistent Layer 3 network boundaries from forming a neighbor relationship. **area number** must also be identical: an OSPF interface belongs to one area, and neighbors exchange LSAs within the context of that area. A common area ID is therefore required before the routers can progress to a full adjacency. Finally, the **hello timer setting** must match. OSPF Hello packets advertise the Hello interval and the associated dead interval; neighbors use these values to coordinate keepalives and failure detection. Matching timing values ensures each router has the same expectation for how frequently Hellos arrive and when a neighbor should be declared unavailable. These compatibility checks occur during neighbor discovery, before the routers synchronize link-state databases. Cisco documents the OSPF Hello packet fields and neighbor-establishment requirements in its [OSPF Neighbor Relationship troubleshooting guide](#) and [OSPF configuration guide](#).

QUESTION NO: 63

What are two features of 6RD IPv6 transition mechanism? (Choose two.)

- A. It inserts IPv4 bits into an IPv6 delegated prefix.
- B. It uses a native IPv6-routed network between CE routers and the BR router.
- C. It allows dynamic 1:N translation of IPv6 address.
- D. It uses stateful automatic 6to4 tunnels between CE routers and the BR router.
- E. It uses stateless automatic 6to4 tunnels between CE routers and the BR router.

ANSWER: A E

Explanation:

6rd (IPv6 Rapid Deployment) provides IPv6 connectivity over an IPv4 service-provider infrastructure by using stateless, automatically derived tunnels between customer-edge routers and 6rd border relays. A customer-edge router derives its 6rd IPv6 prefix from the provider's assigned 6rd prefix and embeds all or part of its IPv4 address in that prefix. This deterministic construction enables the border relay to identify the appropriate IPv4 tunnel endpoint directly from the destination IPv6 address, without maintaining per-customer tunnel state. Therefore, *It inserts IPv4 bits into an IPv6 delegated prefix.* is a defining feature of 6rd.

The provider IPv4 network remains the transport underlay, while IPv6 packets are encapsulated in IPv4 for transport between the customer edge and the 6rd border relay. The tunnel behavior is automatic and stateless, based on the address-derived mapping rather than individually configured tunnel sessions. Consequently, *It uses stateless automatic 6to4 tunnels between CE routers and the BR router.* describes the tunnel model used by 6rd, although 6rd uses a provider-specific prefix and relay design rather than the public 6to4 anycast infrastructure. RFC 5969 defines the 6rd mechanism and its address mapping behavior: [RFC 5969](#). Cisco's implementation guidance is available in the [Cisco IOS IPv6 6rd configuration documentation](#).

QUESTION NO: 64

```
flow record flow
match ipv4 source address
match ipv4 destination address
match ipv4 protocol
match transport source-port
match transport destination-port
collect interface input
collect interface output
collect ipv4 tos
collect transport tcp flags
collect flow direction

flow exporter flow_exporter
destination 172.21.4.1
source loopback 1
transport udp 62123
```

Refer to the exhibit. The network administrator is enabling network visibility on the edge router. The flow exporter must be leveraging an **RFC 7011-based protocol**, and data streams must be sent **every 3 seconds**. Which two commands must the administrator apply to complete the task?

- A. export-protocol netflow-v5
- B. template data timeout 3
- C. export-protocol netflow-v9
- D. export-protocol ipfix
- E. template data timeout 180

ANSWER: B D

Explanation:

`export-protocol ipfix` is required because IP Flow Information Export (IPFIX) is the protocol specified by RFC 7011. IPFIX standardizes the export of flow information from a metering/exporting device to a collecting process. In Cisco Flexible NetFlow, the flow exporter's export protocol must therefore be configured as IPFIX when the requirement explicitly calls for an RFC 7011-based protocol. RFC 7011 also defines the template-based structure used by IPFIX, allowing collectors to interpret the fields contained in exported flow records.

`template data timeout 3` configures the interval at which template data is resent by the flow exporter, expressed in seconds. A three-second timeout meets the stated requirement for the relevant export data stream to be sent every three seconds. Template exports are important because they describe the fields and record layout that the collector uses to decode the IPFIX flow data it receives. Together, these commands configure an IPFIX-capable exporter and set its template-data refresh interval to three seconds. See [RFC 7011: IP Flow Information Export \(IPFIX\) Protocol Specification](#) and Cisco's [Flexible NetFlow Configuration Guide](#).

QUESTION NO: 65

Which two tasks must you perform when you implement LDP NSF on your network? (Choose two.)

- A. Enable NSF for BGP.
- B. Implement direct connections for LDP peers.
- C. Enable NSF for EIGRP.
- D. Disable Cisco Express Forwarding.
- E. Enable NSF for the link-state routing protocol that is in use on the network.

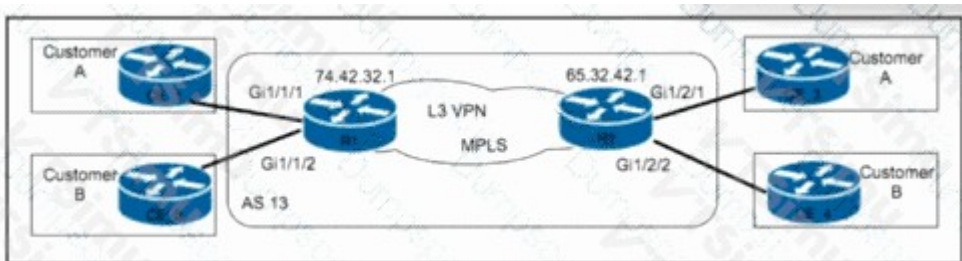
ANSWER: B E

Explanation:

LDP Nonstop Forwarding is intended to preserve MPLS forwarding through a control-plane switchover or routing-process restart. For LDP to provide this protection, the underlying link-state IGP must itself be NSF-capable and enabled for NSF. The IGP retains or recovers the topology and routes needed to support existing label-switched paths while the control plane recovers, minimizing loss of forwarding continuity. Therefore, enabling NSF for the link-state routing protocol in use is a required part of an LDP NSF design.

LDP NSF also requires LDP peers to be directly connected. The feature is designed around preservation of label bindings and forwarding state with directly connected LDP adjacencies during the graceful-restart recovery period. Direct LDP sessions allow the peer relationship and retained forwarding information to be used while LDP reestablishes its control-plane state. Cisco's LDP NSF documentation describes the dependency on NSF-capable IGP operation and the direct-connect requirement for peers. The protocol behavior and restart procedures are also defined by the LDP graceful-restart specification. See Cisco's [MPLS LDP Nonstop Forwarding configuration guide](#) and [RFC 3478](#).

QUESTION NO: 66



Refer to the exhibit. A Tier-2 ISP provides MPLS Layer 3 VPNs to customers throughout the region. BGP and OSPF are used as routing protocols in the backbone, and all routers are running Cisco IOS-XE software. Customer A has requested a private connection between two of its offices. The company has plans to interconnect more offices across the country soon, so the solution should allow for the addition of new sites without reconfiguring the existing sites. The customer is using an IPv6 addressing scheme. R1 has been preconfigured with vrf definition Customer_A and assigned rd 1441:3421.

Which two tasks should the engineer perform on R1 to configure the private connection? (Choose two.)

- A. Configure address-family unicast ipv6 Customer_A under the BGP configuration.
- B. Configure export route-target 1441:3421 under the VRF configuration.
- C. Configure address-family ipv6 unicast under the VRF configuration.
- D. Configure route-target both 1441:3421 under the address family configuration.
- E. Configure address-family unicast ipv6 vrf Customer_A under the OSPF configuration.

ANSWER: C D

Explanation:

An MPLS Layer 3 VPN carrying Customer A's IPv6 routes requires an IPv6 address-family within the Customer_A VRF. The VRF address family provides the per-VRF IPv6 routing context in which customer-facing interfaces, IPv6 routes, and VPN route-target policy are applied. Because the route distinguisher is already configured, defining the IPv6 address family completes the required protocol-specific VRF structure.

Configuring *route-target both 1441:3421* in that address family applies both import and export route-targets. Exporting the route target tags Customer_A IPv6 VPN routes when they are advertised through the provider's MP-BGP VPNv6 control plane. Importing the same route target installs matching Customer_A VPNv6 routes received from other provider edge routers into this VRF. This shared route-target model lets additional customer sites join the same VPN by using the same route-target policy, without requiring changes to the existing sites' routing configuration. Route distinguishers provide uniqueness for VPN prefixes, while route targets control VPN membership and route distribution. Cisco's MPLS VPN configuration guidance describes configuring route targets under the VRF address family and using import/export policies for VPN route exchange. See [Cisco MPLS Layer 3 VPN Overview](#) and [Cisco IPv6 Multiprotocol BGP Configuration Guide](#).

QUESTION NO: 67

```
R1:
interface FastEthernet0/0
ip address 10.1.12.1 255.255.255.0
duplex full
end
!
!
!
R1(config)#interface FastEthernet0/0
R1(config-if)#ospfv3 1 area 1 ipv4
% IPv6 routing not enabled
```

Refer to the exhibit. A network engineer is implementing an OSPF configuration. Based on the output, which statement is true?

- A. OSPFv3 does not run for IPv4 on FastEthernet0/0 until IPv6 routing is enabled on the router and IPv6 is enabled on interface FastEthernet0/0.
- B. In the ospfv3 1 area 1 ipv4 command, area 0 must be configured instead of area 1.
- C. OSPFv3 cannot be configured for IPv4; OSPFv3 works only for IPv6.
- D. "IPv6 routing not enabled" is just an informational message and OSPFv3 runs for IPv4 on interface FastEthernet0/0 anyway.

ANSWER: A

Explanation:

OSPFv3 can carry IPv4 routing information when it is configured with the IPv4 address family, as shown by the interface-level `ospfv3 1 ipv4 area 1` syntax. However, OSPFv3 uses IPv6 for its neighbor communication and protocol transport, including IPv6 link-local addresses on the participating link. Therefore, the router must have IPv6 routing enabled globally with `ipv6 unicast-routing`, and IPv6 must be enabled on FastEthernet0/0 so the interface can support the OSPFv3 IPv6 transport required by the process. The displayed "IPv6 routing not enabled" message indicates that this prerequisite has not yet been met; it is not merely informational for a working IPv4 OSPFv3 deployment. Once IPv6 routing is enabled globally and IPv6 is enabled on the interface, the OSPFv3 IPv4 address-family configuration can establish adjacencies and advertise IPv4 prefixes in the configured area. Cisco documents that OSPFv3 supports separate address families, including IPv4, while retaining IPv6-based OSPFv3 operation. See Cisco's [OSPFv3 configuration guide](#) and [IPv6 configuration overview](#).

QUESTION NO: 68

What are two features of stateful NAT64?

- A. It provides 1: N translations, so it supports an unlimited number of endpoints
- B. It provides 1:1 translation so it supports a limited number of end points
- C. It requires the ipv6 hosts to use either DHCPv6 based address assignments or manual address assignments
- D. It uses address overloading
- E. It requires IPv4 translatable IPv6 address assignments

ANSWER: A D

Explanation:

Stateful NAT64 maintains per-flow translation state so that IPv6-only clients can communicate with IPv4-only destinations. It provides a many-to-one translation model: multiple IPv6 endpoints can use a smaller pool of IPv4 addresses, with each active session distinguished by transport-layer information such as protocol and port numbers. This substantially increases the number of IPv6 clients that can access IPv4 services compared with a fixed one-to-one address mapping. The practical scale is determined by the available IPv4 address pool, ports, and device resources, rather than being literally infinite.

Address overloading is the mechanism that enables this sharing. Similar to IPv4 PAT, the NAT64 device maps multiple IPv6 source sessions to an IPv4 address and allocates unique translated source ports as needed. It records these bindings in its state table so that return traffic from the IPv4 network is mapped back to the originating IPv6 host and flow. Cisco documents stateful NAT64 as supporting many IPv6 hosts through an IPv4 address pool and port translation. See Cisco's [NAT64 configuration guide](#) and the IETF's [Stateful NAT64 RFC](#).