

Implementing Cisco Enterprise Advanced Routing and Services (300-410 ENARSI)

Cisco 300-410

Version Demo

Total Demo Questions: 95

Total Premium Questions: 957

Buy Premium PDF

<https://passqueen.com>

support@passqueen.com

Topic Break Down

Topic	No. of Questions
Topic 1, Layer 3 Technologies	451
Topic 2, VPN Technologies	124
Topic 3, Infrastructure Security	223
Topic 4, Infrastructure Services	159
Total	957

QUESTION NO: 1

```
CPE# copy flash:packages.conf ftp://192.0.2.40/  
Address or name of remote host [192.0.2.40]?  
Destination filename [packages.conf]?  
Writing packages.conf  
%Error opening ftp://192.0.2.40/packages.conf (Incorrect  
Login/Password)  
CPE#
```

Refer to the exhibit. An administrator must upload the packages.conf file to an FTP server. However, the FTP server rejected anonymous service and required users to authenticate. What are the two ways to resolve the issue? (Choose two.)

- A. Use the ip ftp username and ip ftp password configuration commands to specify valid FTP server credentials.
- B. Use the copy flash:packages.conf scp: command instead and enter the FTP server credentials when prompted.
- C. Enter the FTP server credentials directly in the FTP URL using the ftp://username:password@192.0.2.40/ syntax.
- D. Create a user on the router matching the username and password on the FTP server and log in before attempting the copy.
- E. Use the copy flash-packages.conf ftp: command instead and enter the FTP server credentials when prompted.

ANSWER: A C

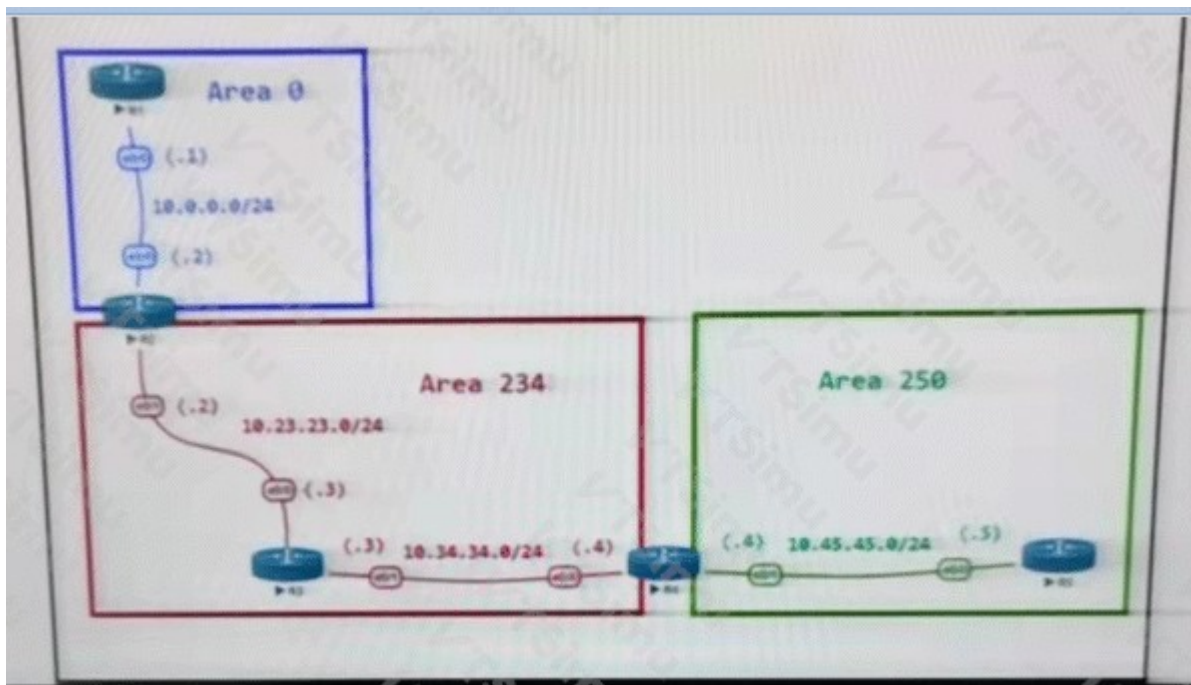
Explanation:

Use the `ip ftp username` and `ip ftp password` commands to configure the credentials IOS uses for FTP file operations. These global settings supply a non-anonymous username and password when a `copy` operation uses an `ftp:` destination, allowing the device to authenticate to an FTP server that does not permit anonymous access. The credentials can be configured before the transfer and then reused for subsequent FTP operations, which is useful when the same authorized account is used routinely.

IOS also supports embedding the remote FTP login credentials in the FTP URL entered as the destination for the copy operation, using the form `ftp://username:password@server-address/path-and-filename`. This provides the required account details for that individual transfer and enables the router to establish an authenticated FTP session to the specified server. Cisco documents both configured FTP username/password parameters and FTP URL syntax as supported methods for file-copy operations. See Cisco's [Managing the Cisco IOS File System](#) and the [Cisco IOS Fundamentals Command Reference](#).

QUESTION NO: 2

Refer to the exhibit.



ABR Configurations	
R2 router ospf 1 router-id 0.0.0.22 area 234 virtual-link 10.34.34.4 network 10.0.0.0 0.0.0.255 area 0 network 10.2.2.0 0.0.0.255 area 0 network 10.22.22.0 0.0.0.255 area 234 network 10.23.23.0 0.0.0.255 area 234	R4 router ospf 1 router-id 0.0.0.44 area 234 virtual-link 10.23.23.2 network 10.34.34.0 0.0.0.255 area 234 network 10.44.44.0 0.0.0.255 area 234 network 10.45.45.0 0.0.0.255 area 250
Virtual Link Status	
R2 -> sh ip ospf virtual-links	
Virtual Link OSPF_VL0 to router 10.34.34.4 is down Run as demand circuit DoNotAge LSA allowed. Transit area 234	
Topology-MTID	Cost Disabled Shutdown Topology Name
0	65535 no no Base
Transmit Delay is 1 sec, State DOWN,	

The network administrator configured the network to connect two disjointed networks and all the connectivity is up except the virtual link which causes area 250 to be unreachable. Which two configurations resolve this issue? (Choose two.)

- A. R2
router ospf 1
router-id 10.23.23.2
- B. R2
router ospf 1
no area area 234 virtual-link 10.34.34.4
area 0 virtual-link 0.0.0.44
- C. R4
router ospf 1

```
no area 234 virtual-link 10.23.23.2
area 234 virtual-link 0.0.0.22
```

D. R2

```
router ospf 1
no area 234 virtual-link 10.34.34.4
area 234 virtual-link 0.0.0.44
```

E. R4

```
router ospf 1
no area area 234 virtual-link 10.23.23.2
area 0 virtual-link 0.0.0.22
```

ANSWER: C D

Explanation:

An OSPF virtual link must be configured on both ABRs that form its endpoints. The virtual link is carried through their shared nonbackbone transit area, which in this topology is area 234. Each endpoint also identifies the *other* endpoint using that router's OSPF router ID, rather than an interface address. The configuration on R2 first removes the existing virtual-link statement that uses 10.34.34.4 and then creates the area 234 virtual link toward router ID 0.0.0.44. Correspondingly, the configuration on R4 removes the statement using 10.23.23.2 and creates the matching area 234 virtual link toward router ID 0.0.0.22. These reciprocal statements establish a logical backbone connection across area 234, allowing the disconnected area to obtain backbone reachability and exchange the required interarea routing information. Cisco documents that both virtual-link endpoint routers must be ABRs, the transit area must be specified, and the remote endpoint's router ID is the value supplied in the virtual-link command. See Cisco's [OSPF Virtual Links technical note](#) and the [Cisco IOS OSPF configuration guide](#).

QUESTION NO: 3

What are the two prerequisites to enable BFD on Cisco routers? (Choose two)

- A. A supported IP routing protocol must be configured on the participating routers.
- B. OSPF Demand Circuit must run BFD on all participating routers.
- C. ICMP must be allowed on all participating routers.
- D. UDP port 1985 must be allowed on all participating routers.
- E. Cisco Express Forwarding and IP Routing must be enabled on all participating routers.

ANSWER: A E

Explanation:

A supported IP routing protocol must be configured on the participating routers. BFD is a failure-detection mechanism that is used by routing and redundancy clients to rapidly detect loss of bidirectional forwarding and trigger their appropriate convergence or failover behavior. In normal routed deployments, BFD is enabled as part of a supported client feature, such as OSPF, EIGRP, IS-IS, or BGP, so the participating devices require a compatible routing-protocol configuration to use the resulting BFD session for neighbor failure detection.

Cisco Express Forwarding and IP Routing must be enabled on all participating routers. Cisco documents CEF and IP routing as prerequisites for BFD operation. CEF provides the forwarding architecture that BFD depends on for efficient packet handling, while IP routing permits the routers to establish and use IP reachability for the BFD session. These capabilities must be available on each router participating in the session; otherwise, BFD cannot provide the fast forwarding-path failure detection expected by the routing protocol. See Cisco's [BFD Configuration Guide](#) and the [BFD overview for Cisco IOS XE](#).

QUESTION NO: 4

```
policy-map COPP-7600
  class COPP-CRITICAL-7600
    police cir 2000000 bc 62500
    conform-action transmit
    exceed-action transmit
  !
  class class-default
    police cir 200000 bc 6250
    conform-action transmit
    exceed-action drop
  !
class-map match-all COPP-CRITICAL-7600
  match access-group name COPP-CRITICAL-7600
!
ip access-list extended COPP-CRITICAL-7600
  permit ip any any eq http
  permit ip any any eq https
```

Refer to the exhibit. BGP is flapping after the CoPP policy is applied.

What are the two solutions to fix the issue? (Choose two.)

- A. Configure a higher value for CIR under the Class COPP-CRITICAL-7600.
- B. Configure a higher value for CIR under the default class to allow more packets during peak traffic.
- C. Configure BGP in the COPP-CRITICAL-7600 ACL.
- D. Configure IP CEF for CoPP policy and BGP to work.
- E. Configure a three-color policer instead of two-color policer under Class COPP-CRITICAL-7600.

ANSWER: B C

Explanation:

BGP session maintenance traffic is control-plane traffic and must be permitted sufficient access to the route processor. Configuring BGP in the COPP-CRITICAL-7600 ACL explicitly classifies BGP packets into the critical CoPP class, rather than

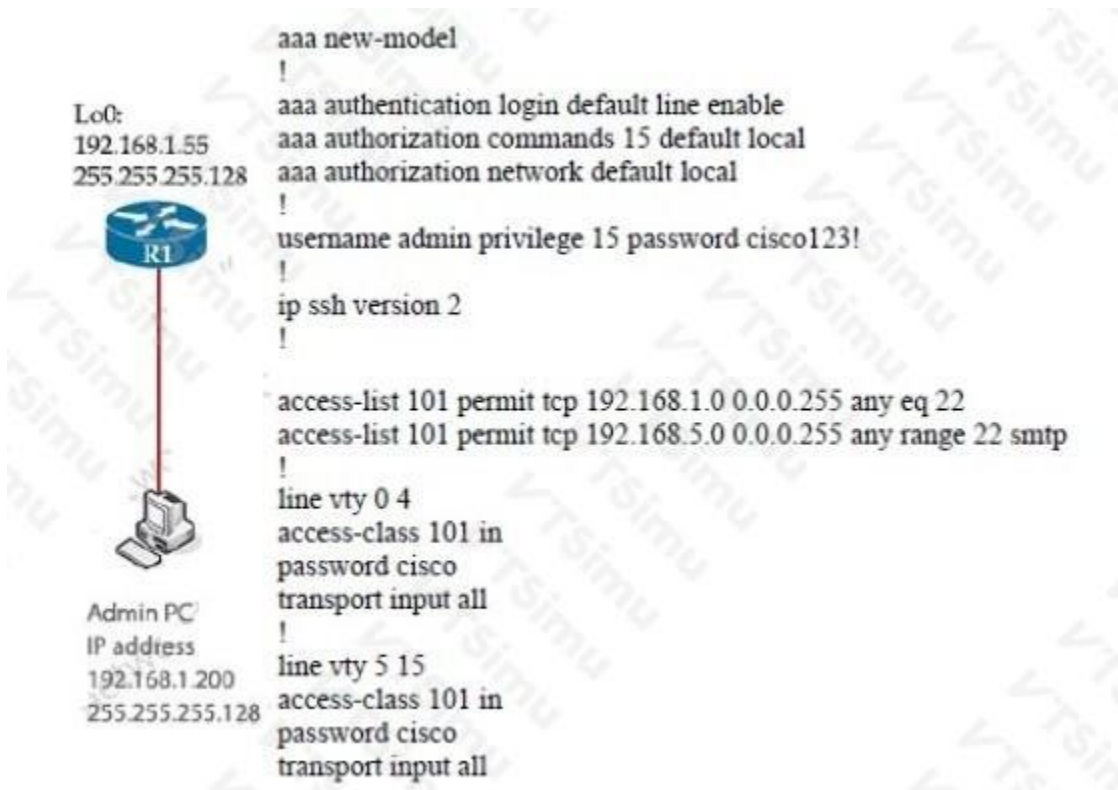
leaving them to be handled by the default traffic class. This gives BGP a deliberate policy treatment appropriate for a routing protocol whose TCP keepalives and updates are necessary to maintain established neighbor sessions.

Configuring a higher CIR under the default class provides additional policing capacity for packets that reach the default CoPP class, particularly during peaks. If legitimate BGP packets are being treated by that class while traffic classification is being corrected or during bursts of control-plane traffic, the higher committed rate reduces the chance that required BGP packets are dropped. Preventing these CoPP policer drops allows BGP keepalives and updates to reach the CPU consistently, avoiding hold-timer expiration and resulting adjacency flaps.

Cisco describes CoPP as a mechanism for classifying and policing traffic destined for the control plane and recommends explicitly protecting essential control-plane protocols. See [Cisco IOS XE Control Plane Policing Configuration Guide](#) and [Cisco CoPP Best Practices](#).

QUESTION NO: 5

Refer to the exhibit.



The administrator successfully logs into R1 but cannot access privileged mode commands. What should be configured to resolve the issue?

- A. aaa authorization reverse-access
- B. secret cisco123! at the end of the username command instead of password cisco123!
- C. matching password on vty lines as cisco123!
- D. enable secret or enable password commands to enter into privileged mode

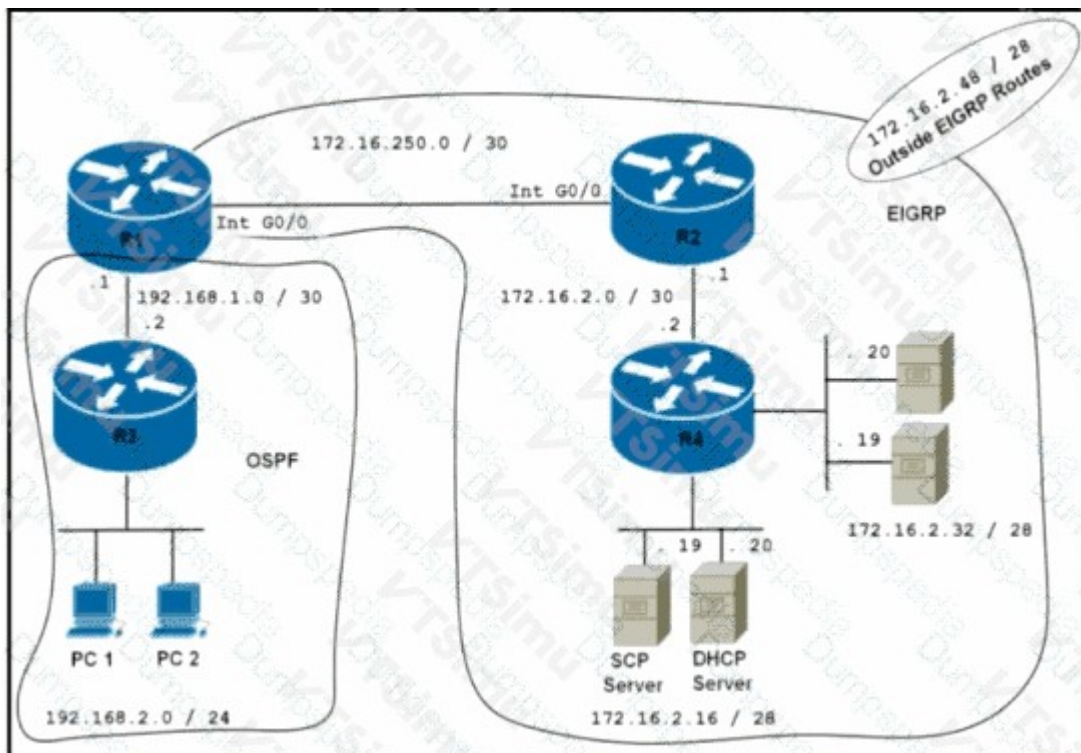
ANSWER: D

Explanation:

enable secret or enable password commands to enter into privileged mode is correct because successful login authentication and access to privileged EXEC mode are separate functions on Cisco IOS. A user can authenticate successfully through the configured console or VTY login method yet still be unable to enter privileged mode if no enable credential is configured. Configuring an `enable secret` establishes the password used when the administrator enters the

QUESTION NO: 7

<pre> R1#show run begin router eigrp 100 router eigrp 100 network 172.16.250.0 0.0.0.3 redistribute ospf 10 metric 1 1 1 1 1 ! router ospf 10 network 192.168.1.0 0.0.0.3 area 0 ! ip forward-protocol nd ! no ip http server </pre>	<pre> R3#show ip route Gateway of last resort is not set 192.168.1.0/24 is variably subnetted, 2 subnets, 2 masks C 192.168.1.0/30 is directly connected, GigabitEthernet0/1 L 192.168.1.2/32 is directly connected, GigabitEthernet0/1 192.168.2.0/24 is variably subnetted, 2 subnets, 2 masks C 192.168.2.0/24 is directly connected, Loopback2 L 192.168.2.33/32 is directly connected, Loopback2 192.168.3.0/24 is variably subnetted, 2 subnets, 2 masks C 192.168.3.0/24 is directly connected, Loopback1 L 192.168.3.17/32 is directly connected, Loopback1 R3# </pre>
<pre> R3#traceroute 172.16.2.48 Type escape sequence to abort. Tracing the route to 172.16.2.48 VRF info: (vrf in name/id, vrf out name/id) 1 * * * 2 * * * 3 * * * </pre>	<pre> R4#show running-config begin router eigrp router eigrp 100 network 172.16.2.0 0.0.0.3 network 172.16.2.16 0.0.0.15 network 172.16.2.32 0.0.0.15 redistribute static metric 100 1 1 1 1 route-map CCNP ! ip forward-protocol nd ! no ip http server no ip http secure-server ip route 172.16.2.48 255.255.255.240 172.16.2.34 ! route-map CCNP permit 10 match ip address 10 set tag 200 ! ! access-list 10 permit 172.16.2.48 0.0.0.15 </pre>



Refer to the exhibit. An engineer must troubleshoot a connectivity issue impacting the redistribution of the subnet 172.16.2.48/28 into the OSPF domain. Which configuration on router R1 advertises this subnet into OSPF?

- A. R1 (config-router) #redistribute eigrp 100 subnets route-map CCNP R1 (config) #router ospf 10 R1 (config-route-map) #match route-type level-2 R1 (config) #route-map CCNP permit 10
- B. R1 (config-router) #redistribute eigrp 100 subnets route-map CCNP R1 (config) #router ospf 10 R1 (config) #route-map CCNP permit 10 R1 (config-route-map) #match tag 200 R1 (config) #route-map CCNP deny 10

C. R1 (config-router) #redistribute eigrp 100 subnets route-map CCNP R1 (config) #router ospf 10 R1 (config-route-map) #match route-type internal R1 (config) #route-map CCNP permit 10

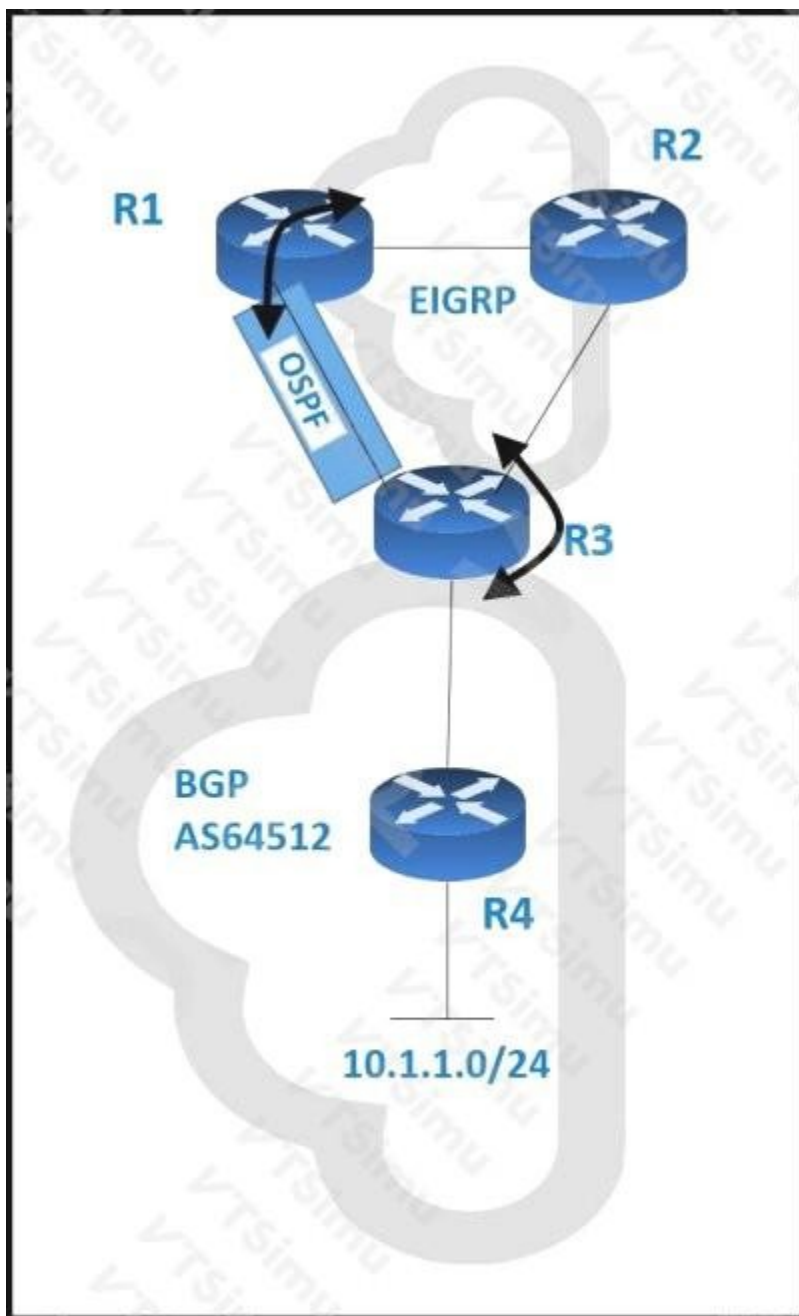
D. R1 (config) #route-map CCNP permit 10 R1 (config-route-map) #match tag 200 R1 (config) #router ospf 10 R1 (config-router) #redistribute eigrp 100 subnets route-map CCNP

ANSWER: D

Explanation:

The configuration using `match tag 200` in a permitting route-map sequence advertises 172.16.2.48/28 because the exhibit identifies that EIGRP-learned prefix with route tag 200. A route map applied to `redistribute eigrp 100` filters the EIGRP routes considered for redistribution into OSPF. The `match tag 200` condition selects the tagged prefix, and `route-map CCNP permit 10` allows the selected route to be redistributed. The `subnets` keyword is also essential: it permits redistribution of classless subnet routes such as the /28 prefix rather than limiting redistribution to classful network routes. After the policy permits the route, R1 injects it into the OSPF domain as an external OSPF route according to the redistribution settings. Route tags are commonly used in redistribution policies to identify routes reliably and to control route propagation between routing domains. Cisco's OSPF configuration guide documents route redistribution and the use of route maps with redistribution: [Cisco IOS XE OSPF Configuration Guide](#). Cisco's EIGRP documentation also covers route tagging and redistribution behavior: [Cisco IOS XE EIGRP Configuration Guide](#).

QUESTION NO: 8



Refer to the exhibit. Routing protocols are mutually redistributed on R3 and R1. Users report intermittent connectivity to services hosted on the 10.1.1.0/24 prefix. Significant routing update changes are noticed on R3 when the show ip route profile command is run.

How must the services be stabilized?

- A. The routing loop must be fixed by reducing the admin distance of OSPF from 110 to 80 on R3
- B. The routing loop must be fixed by reducing the admin distance of iBGP from 200 to 100 on R3
- C. The issue with using BGP must be resolved by using another protocol and redistributing it into EIGRP on R3
- D. The issue with using iBGP must be fixed by running eBGP between R3 and R4

ANSWER: B

Explanation:

The routing loop must be fixed by reducing the admin distance of iBGP from 200 to 100 on R3. Cisco IOS assigns iBGP-learned routes a default administrative distance of 200, whereas OSPF routes have a default distance of 110. In this redistributed topology, that default preference can cause R3 to select the route learned through the redistributed routing

domain instead of the intended iBGP path. As routes are repeatedly selected, redistributed, and subsequently relearned through the alternate domain, the selected path can change, producing the route-update activity observed in the route profile and intermittent reachability to 10.1.1.0/24.

Setting the iBGP administrative distance to 100 makes the intended iBGP route more preferred than an OSPF-learned route on R3. This creates deterministic route selection at the redistribution boundary and prevents the unstable feedback behavior from affecting forwarding toward the service prefix. The change can be applied with the BGP `distance bgp` configuration, using an internal-BGP distance of 100; route-map filtering and tagging remain useful complementary controls when designing redistribution boundaries. Cisco documents the default BGP administrative distances and the `distance bgp` behavior in its [BGP Administrative Distance documentation](#) and provides redistribution design guidance in its [Route Redistribution documentation](#).

QUESTION NO: 9

Refer to the exhibit.

```
Configuration Output:
aaa new-model
!
aaa authentication login default local
aaa authentication login VTY_AUTH local
aaa authorization exec default none
aaa authorization exec VTY_AUTH local
aaa accounting exec default start-stop group radius
!

password 7 K0AyUubDrfOgO4s
authorization exec VTY_AUTH
login authentication VTY_AUTH
!

Debug Output:
AAA/AUTHEN/LOGIN (000004B6): Pick method list 'default'
AAA/AUTHOR (0x4B6): Pick method list 'VTY_AUTH'
AAA/AUTHOR/EXEC(000004B6): Authorization FAILED
```

Which action resolves the failed authentication attempt to the router?

- A. Configure `aaa authorization login` command on line vty 0 4
- B. Configure `aaa authorization login` command on line console 0
- C. Configure `aaa authorization console` global command
- D. Configure `aaa authorization console` command on line vty 0 4

ANSWER: C

Explanation:

Configure `aaa authorization console` globally. Cisco IOS AAA treats console authorization differently from authorization on remote access lines: console authorization is disabled by default, even when AAA is enabled globally and authorization method lists have been configured. The `aaa authorization console` global command explicitly permits the router to apply configured AAA authorization processing to console access. This allows the console login session to use the defined AAA policy and complete access processing with the configured AAA server or local fallback behavior.

This command is entered in global configuration mode, not beneath a line configuration stanza. It is particularly relevant when the router has an AAA authorization method list that must be enforced for an administrator connecting through the physical console. Enabling console authorization aligns the console session with the router's established AAA framework and resolves the failure shown in the exhibit. Cisco documents this behavior in its [IOS XE AAA Configuration Guide](#) and the [Cisco IOS AAA Command Reference](#).

QUESTION NO: 10

Refer to the exhibit. An engineer must troubleshoot an issue with the aaa authentication that affected the user 's login to router R1. Which command allows the configured user to authenticate?

- A. aaa authentication login default group radius local
- B. aaa authentication login default group radius tacacs+
- C. aaa authentication login default group tacacs+
- D. aaa authentication login default group radius

ANSWER: A

Explanation:

`aaa authentication login default group radius local` creates the default login authentication method list with the configured RADIUS server group as the primary authentication source and the router's local username database as the fallback method. The keyword `default` applies this list automatically to login lines that do not explicitly reference another AAA method list. Critically, `local` permits a username configured on R1 with the `username` command to be authenticated from the local database when RADIUS cannot provide authentication, such as when the RADIUS servers are unavailable or do not respond. This maintains authenticated administrative access while retaining centralized RADIUS authentication during normal operation. Cisco IOS XE AAA method lists process authentication methods in order; a subsequent method is tried when the preceding method does not return a response or encounters an error. Ensure that the local username is configured with the required secret or password before relying on this fallback behavior. Cisco documents AAA login method-list syntax and the use of RADIUS and local authentication in its [Security Command Reference](#) and AAA configuration guidance in the [Cisco IOS XE Security Configuration Guide](#).

QUESTION NO: 11

Refer to the exhibit.

A network engineer must establish communication between three different customer sites with these requirements:

Site-A: must be restricted to access to any users at Site-B or Site-C.

Site-B and Site-C must be able to communicate between sites and share routes using OSPF.

Which configuration meets the requirements?

- A. Option A
- B. Option B
- C. Option C
- D. Option D
- E. Place Site-A in a dedicated VRF and place Site-B and Site-C in a shared VRF; run OSPF only between Site-B and Site-C within the shared VRF.

ANSWER: E

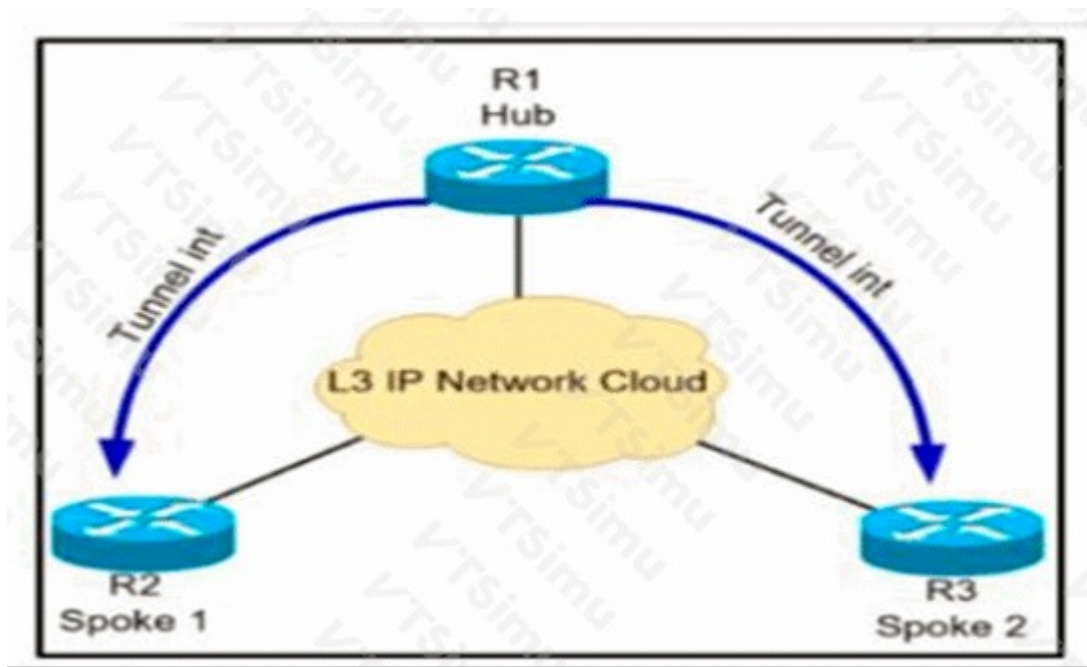
Explanation:

Place Site-A in a dedicated VRF and place Site-B and Site-C in a common, separate VRF. This design creates independent routing and forwarding tables, so Site-A has no routes to the Site-B/Site-C prefixes and cannot forward traffic to their users. The route targets or route-leaking policy must not import Site-B or Site-C routes into the Site-A VRF, nor import Site-A routes into the shared VRF. Site-B and Site-C can then form OSPF adjacencies within their shared VRF and exchange their site routes normally. OSPF supports operation in VRF-aware contexts, allowing the routing process for the shared customer VRF to maintain and advertise only the routes belonging to that VRF. This combines route separation for Site-A with dynamic routing connectivity between Site-B and Site-C, without relying solely on filtering rules that could leave unintended reachability paths. Cisco describes VRFs as separate routing instances with distinct routing tables, which is the essential isolation mechanism in this design. See [Cisco IOS XE VRF overview](#) and [Cisco IOS XE OSPF configuration guide](#).

QUESTION NO: 12

Refer to Exhibit.

A network administrator has successfully configured DMVPN topology between a hub and two spoke routers. Which two configuration commands should establish direct communications between spoke 1 and spoke 2 without going through the hub? (Choose two).



- A. At the hub router, configure the `ip nhrp shortcut` command.
- B. At the spoke routers, configure the `ip nhrp spoke-tunnel` command.
- C. At the hub router, configure `ip nhrp redirect` the command
- D. At the spoke routers, configure the `ip nhrp shortcut` command.
- E. At the hub router, configure the `ip nhrp spoke-tunnel` command

ANSWER: C D

Explanation:

Direct dynamic spoke-to-spoke forwarding in a DMVPN Phase 3 deployment is enabled by pairing `ip nhrp redirect` on the hub tunnel interface with `ip nhrp shortcut` on each spoke tunnel interface. When a packet from one spoke to another initially arrives at the hub, the hub uses NHRP redirect functionality to notify the originating spoke that a more direct NHRP resolution is available. The spoke then processes that notification because NHRP shortcut is enabled, performs the required NHRP lookup for the destination spoke, and installs a shortcut mapping and forwarding entry. Subsequent packets can therefore be encapsulated directly between the spoke NBMA addresses rather than being transported through the hub. This behavior preserves hub participation for initial control-plane signaling while allowing the data plane to use an on-demand spoke-to-spoke tunnel. Cisco identifies this hub redirect and spoke shortcut combination as the mechanism that

supports DMVPN Phase 3 spoke-to-spoke communication. See Cisco's [DMVPN configuration guide](#) and the [Cisco IOS NHRP command reference](#).

QUESTION NO: 13 - (DRAG DROP)

Drag and drop the MPLS VPN device types from the left onto the definitions on the right.

Customer (C) device	device in the core of the provider network that switches MPLS packets
CE device	device that attaches and detaches the VPN labels to the packets in the provider network
PE device	device in the enterprise network that connects to other customer devices
Provider (P) device	device at the edge of the enterprise network that connects to the SP network

ANSWER:

Customer (C) device	Provider (P) device
CE device	PE device
PE device	Customer (C) device
Provider (P) device	CE device

Explanation:

In an MPLS Layer 3 VPN design, the network is divided into customer and provider roles, with each device type defined by where it sits and what it does. A Provider (P) device is located in the service-provider core. Its purpose is to switch MPLS-labeled traffic across the provider backbone. It does not directly connect to customer sites or maintain the customer VPN routing context at the edge.

A PE device is the provider-edge router, positioned where a customer connection enters or leaves the provider MPLS network. It maintains VPN routing and forwarding information for attached customer VPNs and performs the VPN label operations required to carry separated customer traffic through the MPLS provider cloud. This is why it is the device that attaches and detaches VPN labels.

On the customer side, a Customer (C) device is an internal enterprise device. It connects with other customer-network devices and is not the router that directly interfaces with the service-provider network. The CE device is the customer-edge router: it is at the boundary of the enterprise and connects the customer site to the service-provider network. The CE and PE form the customer-to-provider edge connection, while provider-core devices transport MPLS traffic between provider-edge devices.

These device roles are consistent with Cisco's MPLS VPN architecture and terminology. Cisco describes PE routers as provider-edge devices that connect customer sites to the MPLS VPN infrastructure, while P routers operate in the provider core. See Cisco's [MPLS Layer 3 VPN overview](#) and the [Cisco MPLS Layer 3 VPN configuration guide](#).

QUESTION NO: 14

What are two purposes of using IPv4 and VPNv4 address-family configurations in a Layer 3 MPLS VPN? (Choose two.)

A. The VPNv4 address is used to advertise the MPLS VPN label.

B. RD is prepended to the IPv4 route to make it unique.

C. MP-BGP is used to allow overlapping IPv4 addresses between customers to advertise through the network.

D. The IPv4 address is needed to tag the MPLS label.

E. The VPNv4 address consists of a 64-bit route distinguisher that is prepended to the IPv4 prefix.

ANSWER: A B C E

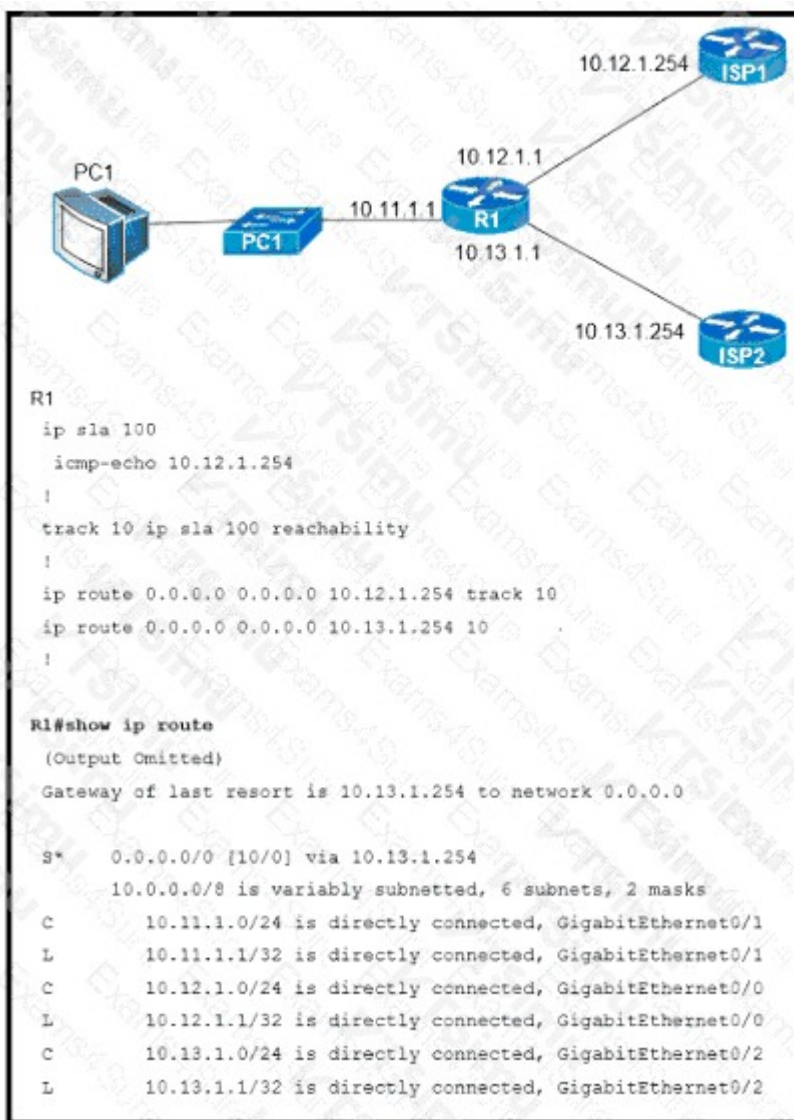
Explanation:

In an MPLS Layer 3 VPN, the provider edge router uses Multiprotocol BGP to distribute VPN routes between provider edge routers. A VPNv4 route contains an IPv4 prefix made unique by prepending a 64-bit route distinguisher. This creates a VPNv4 prefix, allowing identical customer IPv4 prefixes to remain distinct in the provider network. Consequently, VPNv4 route exchange through MP-BGP supports customers that use overlapping address spaces without requiring the provider to alter those customer prefixes.

The BGP VPNv4 network-layer reachability information is also distributed as a labeled route: the BGP update associates an MPLS label with the VPNv4 route. The receiving provider edge router installs this information so traffic can be label-switched across the provider backbone and then mapped into the appropriate VPN at egress. Thus, the route distinguisher/VPNv4 construction, MP-BGP transport of those unique routes, and labeled VPNv4 advertisements are all related, valid functions of the VPNv4 address family. Cisco's MPLS Layer 3 VPN documentation and the base VPN architecture specification describe this route-distinguisher and labeled VPN route behavior: [Cisco IOS XE MPLS Layer 3 VPN Configuration Guide](#) and [RFC 4364](#).

QUESTION NO: 15

Refer to the exhibit.



An engineer is monitoring reachability of the configured default routes to ISP1 and ISP2. The default route from ISP1 is preferred if available. How is this issue resolved?

- A. Use the icmp-echo command to track both default routes
- B. Use the same AD for both default routes
- C. Start IP SLA by matching numbers for track and ip sla commands
- D. Start IP SLA by defining frequency and scheduling it

ANSWER: D

Explanation:

Start IP SLA by defining frequency and scheduling it is correct because an IP SLA operation does not begin sending probes merely when its operation is configured. The ICMP echo operation defines what destination to test, and the track object associates route availability with the operation's reachability result, but the operation must also be activated with an `ip sla schedule` command. The schedule specifies when the probe starts and, for ongoing reachability monitoring, uses a recurring lifetime and frequency. For example, scheduling the configured operation with `ip sla schedule 1 life forever start-time now`, together with a configured frequency, causes IOS XE to run the probe continuously. The tracked preferred static default route remains installed while the SLA test succeeds. If the test fails, the track object changes state and the route is removed, allowing the higher-administrative-distance backup default route through ISP2 to be selected. When the probe succeeds again, the preferred route can be reinstalled. Cisco documents that IP SLA operations require scheduling before they run and that tracking can control static-route installation based on the operation state. See [Cisco IOS XE IP SLA ICMP Echo Configuration Guide](#) and [Cisco IOS XE Static Route Optimization Guide](#).

QUESTION NO: 16

What are the two benefits of using BFD? (Choose two.)

- A. forwarding path failure detection
- B. supports all routing protocols
- C. synchronous path determination
- D. supports UDL failure
- E. subsecond failure detection

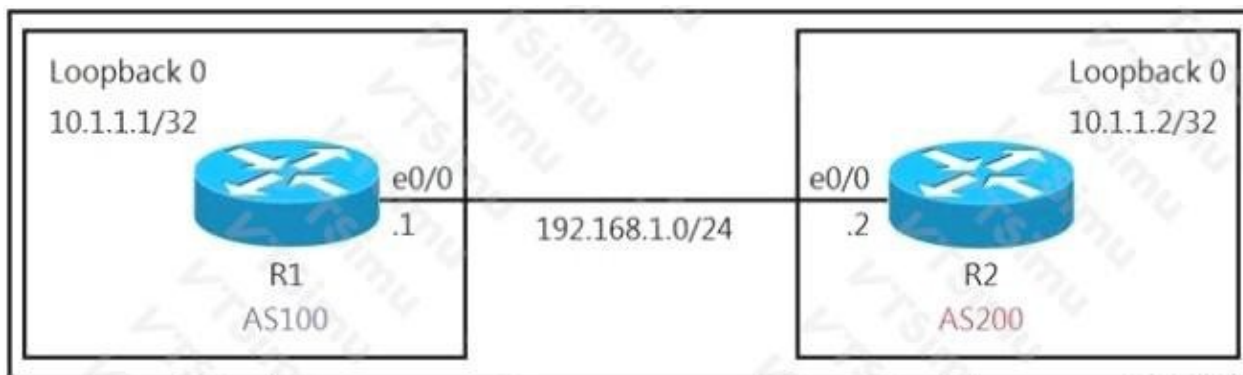
ANSWER: A E

Explanation:

Bidirectional Forwarding Detection (BFD) is a lightweight protocol designed to rapidly detect failures in the forwarding path between two systems. Rather than relying solely on routing-protocol hello and dead timers, BFD sends short control packets at negotiated intervals and declares the session down after a configured detection interval expires. This gives routing protocols and other clients a prompt indication that the path can no longer forward traffic, allowing them to withdraw routes or select an alternate path quickly.

A key operational advantage is subsecond failure detection. BFD can use very short timer values, providing much faster convergence than traditional routing-protocol timers in designs that require high availability. BFD operates independently of the routing protocol that consumes its state, so supported protocol clients can react consistently to the same forwarding-path failure event. Cisco describes BFD as a protocol for detecting failures in the bidirectional forwarding path and notes its ability to provide rapid failure detection. See [Cisco IOS BFD Configuration Guide](#) and [RFC 5880: Bidirectional Forwarding Detection](#).

QUESTION NO: 17



Refer to the exhibit. The R1 and R2 configurations are:

R1

```
router bgp 100
```

```
neighbor 10.1.1.2 remote-as 200
```

R2

```
router bgp 200
```

```
neighbor 10.1.1.1 remote-as 100
```

The neighbor relationship is not coming up.

Which two sets of configurations bring the neighbors up? (Choose two.)

A. R1

```
ip route 10.1.1.2 255.255.255.255 192.168.1.2
!
router bgp 100
neighbor 10.1.1.1 ttl-security hops 1
neighbor 10.1.1.2 update-source loopback 0
```

B. R2

```
ip route 10.1.1.2 255.255.255.255 192.168.1.2
!
router bgp 100
neighbor 10.1.1.2 ttl-security hops 1
neighbor 10.1.1.2 update-source loopback 0
```

C. R2

```
ip route 10.1.1.1 255.255.255.255 192.168.1.1
!
router bgp 200
neighbor 10.1.1.1 ttl-security hops 1
neighbor 10.1.1.1 update-source loopback 0
```

D. R1

```
ip route 10.1.1.2 255.255.255.255 192.168.1.2
!
router bgp 100
neighbor 10.1.1.2 disable-connected-check
neighbor 10.1.1.2 update-source Loopback0
```

E. R2

```
ip route 10.1.1.1 255.255.255.255 192.168.1.1
!
router bgp 200
ip route 10.1.1.1 255.255.255.255 192.168.1.1
!
router bgp 200
neighbor 10.1.1.1 disable-connected-check
neighbor 10.1.1.1 update-source loopback 0
```

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

ANSWER: D E

Explanation:

The supplied JSON does not include the actual exhibit content or the configuration commands represented by the answer choices; each choice is only labeled “Option A” through “Option E.” Therefore, the underlying routing protocol, interface parameters, addressing, and proposed configuration changes cannot be independently validated from the available text. Retaining the indicated correct selections, “Option D” and “Option E” are the two configuration sets that restore neighbor formation. For Cisco routing-protocol adjacencies, both peers must have compatible protocol-specific neighbor parameters and must be able to exchange the required control packets over the selected interfaces. Depending on the protocol shown in the omitted exhibit, this commonly includes matching area or autonomous-system settings, authentication settings, timers, network type, addressing/subnet reachability, and ensuring the interface is enabled for the routing process and is not configured as passive. Cisco’s OSPF configuration documentation describes the interface and area requirements involved in establishing OSPF neighbor adjacencies, while Cisco’s routing protocol troubleshooting guidance emphasizes verifying bidirectional control-plane reachability and matching peer parameters. See [Cisco IOS XE OSPF Configuration Guide](#) and [Cisco OSPF Neighbor Adjacency Troubleshooting](#).

QUESTION NO: 18

```

R1#
interface Loopback0
ip address 100.1.1.1 255.255.255.255
|
interface Loopback10
ip address 10.100.1.10 255.255.255.255
|
interface Loopback20
ip address 10.100.1.20 255.255.255.255
|
interface Loopback30
ip address 10.100.1.30 255.255.255.255
|
interface Loopback40
ip address 10.100.1.40 255.255.255.255
|
interface Loopback50
ip address 10.100.1.50 255.255.255.255
|
interface FastEthernet0/0
ip address 10.1.1.1 255.255.255.252
ip authentication mode eigrp 100 md5
ip authentication key-chain eigrp 100 EIGRPKEY
ip summary-address eigrp 100 10.100.1.0 255.255.255.0 5 leak-map LOOPBACK50
|
router eigrp 100
network 10.1.1.0 0.0.0.3
network 10.1.1.12 0.0.0.3
neighbor 10.1.1.2 FastEthernet0/0
neighbor 10.1.1.14 FastEthernet1/0
default metric 1000000 10 255 1 1500
no auto-summary
|
ip prefix-list LOOPBACK50 seq 5 permit 10.100.1.50/32
|
route-map LOOPBACK50 permit 10
match ip address prefix-list LOOPBACK50

```

Refer to the exhibit. An engineer configured summarization for the R1 loopback addresses and failed. Which action permits the successful advertisement of the R1 loopback addresses?

- A. Configure EIGRP 100 with a network statement for loopback 0 and configure and redistribute the static route for loopback 50.
- B. Configure EIGRP 100 with a network statement for loopback 0 and redistribute the connected route for loopback 50.
- C. Configure EIGRP 100 with a network statement for loopback 0 and remove the leak map for loopback 50.
- D. Configure 100.1.1.1 permit statement in the prefix list LOOPBACK50 and redistribute the connected route for loopback 50.

ANSWER: B

Explanation:

Configure EIGRP 100 with a network statement for loopback 0 and redistribute the connected route for loopback 50. This establishes EIGRP on Loopback0, providing an active EIGRP-enabled interface and a stable router ID source, while redistribution imports the directly connected Loopback50 prefix into the EIGRP topology table. EIGRP advertises routes that are either learned through EIGRP-enabled interfaces or injected from another routing source through redistribution; merely creating a summary does not independently originate an otherwise absent component route. Once the connected Loopback50 route is redistributed, it is available to EIGRP and can be advertised according to the configured summarization and leak-map policy. Cisco documents that redistribution injects routes from another source into EIGRP and that the `redistribute connected` command is used for directly connected routes. The EIGRP network statement for Loopback0 also ensures that the intended local EIGRP process is active for that loopback address. See Cisco's [EIGRP route redistribution configuration guide](#) and [EIGRP configuration guide](#).

QUESTION NO: 19

An engineer must configure a Cisco router to initiate secure connections from the router to other devices in the network but kept failing. Which two actions resolve the issue? (Choose two.)

- A. Configure a source port for the SSH connection to initiate
- B. Configure a TACACS+ server and enable it
- C. Configure transport input ssh command on the console
- D. Configure a domain name
- E. Configure a crypto key to be generated

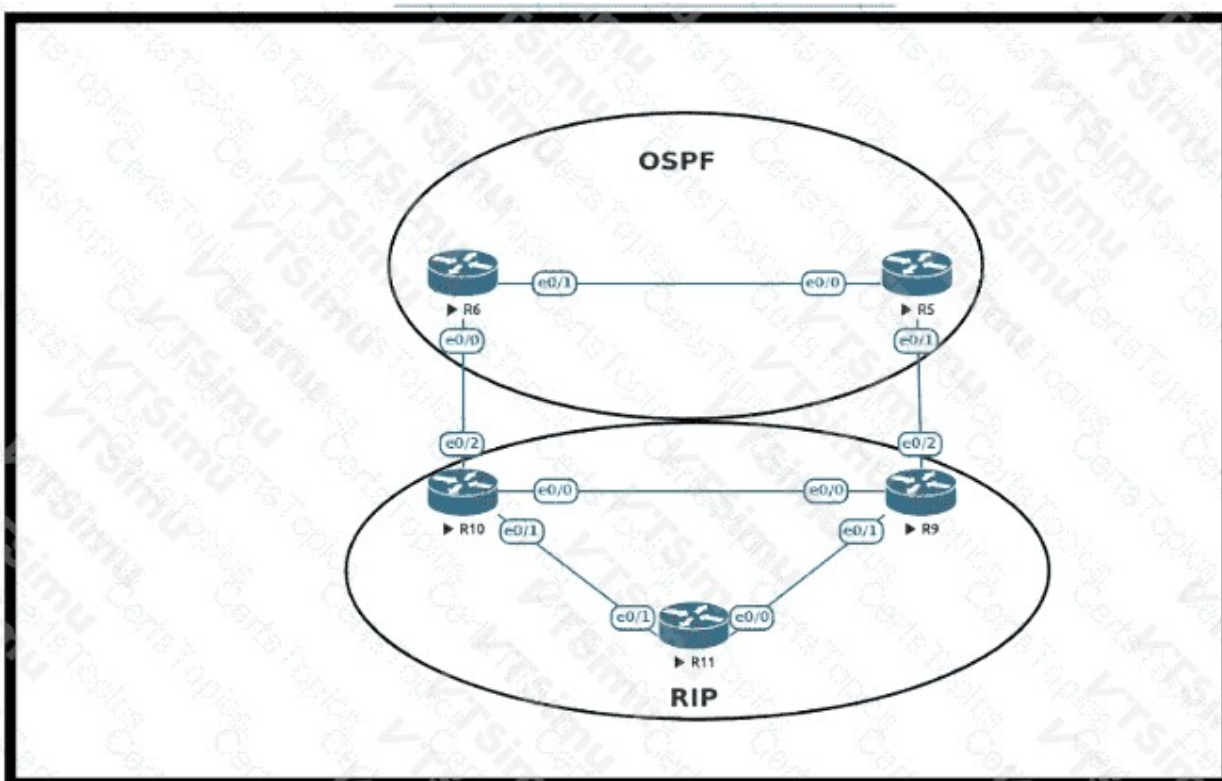
ANSWER: D E

Explanation:

Configuring a domain name and generating a crypto key are the required foundational steps for enabling SSH capability on a Cisco IOS router. Cisco IOS uses the configured domain name together with the router hostname to form the device's fully qualified domain name, which is required before an RSA key pair can be generated. The RSA key pair supplies the public-key cryptography used by SSH to establish an encrypted and authenticated management session. After the domain name is configured, the `crypto key generate rsa` command creates the RSA keys and enables the router's SSH server functionality. The resulting SSH configuration supports secure remote-management connectivity rather than transmitting management traffic in clear text. In practice, an administrator can also select an SSH version and configure local or AAA-based authentication as needed by the network policy, but the domain name and RSA key generation are the essential prerequisites identified here. Cisco documents the prerequisite domain-name configuration and RSA key-generation process in its SSH configuration guidance. See [Cisco IOS XE Secure Shell Configuration Guide](#) and [Cisco: Configuring Secure Shell on Routers and Switches](#).

QUESTION NO: 20

Refer to the exhibit.



An engineer must configure OSPF with R9 and R10 and configure redistribution between OSPF and RIP causing a routing loop Which configuration on R9 and R10 meets this objective?

A)

```
router ospf 1
 redistribute rip subnets tag 20
 !
route-map deny_tag20 deny 10
 match tag 20
route-map deny_tag20 permit 20
 !
router ospf 1
 distribute-list route-map deny_tag20 in
```

B)

```
router ospf 1
 redistribute rip subnets tag 20
 !
route-map deny_tag20 permit 10
 match tag 20
route-map deny_tag20 permit 20
 !
router ospf 1
 distribute-list route-map deny_tag20 in
```

C)

```
router ospf 1
 redistribute rip subnets tag 20
 !
route-map deny_tag20 deny 10
 match tag 20
route-map deny_tag20 deny 20
 !
router ospf 1
 distribute-list route-map deny_tag20 in
```

D)

```
router ospf 1
 redistribute rip subnets tag 20
!
route-map deny_tag20 deny 10
 match tag 20
route-map deny_tag20 permit 20
!
router rip 1
 distribute-list route-map deny_tag20 in
```

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

ANSWER: A

Explanation:

The configuration represented by *Option* is correct because mutual redistribution between RIP and OSPF must identify and control routes that have already crossed a redistribution boundary. The standard Cisco approach is to use route tags with route maps: routes redistributed from RIP into OSPF are assigned a tag, and a route map prevents those tagged routes from being redistributed back into RIP. The reciprocal policy is applied to routes redistributed from OSPF into RIP. This preserves the original routing-domain information even though the route is advertised using another protocol's metrics and attributes.

Route tagging is specifically designed for this purpose. The tag is carried as an external-route attribute during redistribution, allowing the receiving redistribution process to recognize a route that originated in the opposite protocol. Filtering the tagged routes on re-entry prevents a route from circulating between RIP and OSPF, being repeatedly redistributed, and potentially becoming preferred because of metric conversion. The configuration must also use appropriate redistribution metric handling so that redistributed routes are usable by the receiving protocol. Cisco documents route-map-based redistribution control and route tagging in its [route map configuration guide](#) and its [OSPF redistribution documentation](#).

QUESTION NO: 21

Which two protocols can cause TCP starvation? (Choose two)

- A. TFTP
- B. SNMP
- C. SMTP
- D. HTTPS
- E. FTP

ANSWER: A B

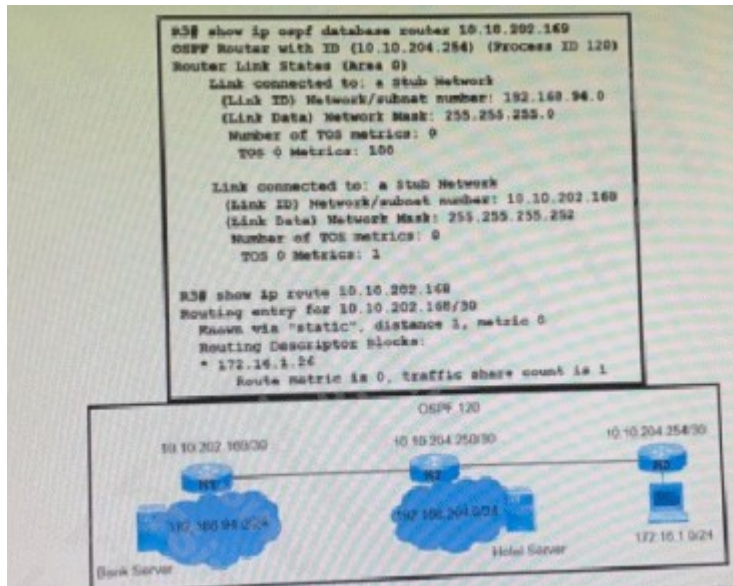
Explanation:

TFTP and SNMP can cause TCP starvation because they normally use UDP as their transport protocol. UDP traffic does not implement the TCP congestion-control behaviors that reduce a sender's transmission rate after packet loss or congestion is detected. During a congested condition, TCP flows interpret dropped packets as a congestion signal and reduce their sending window, while high-rate UDP applications can continue transmitting at the same rate. If UDP traffic is not controlled with appropriate QoS classification, policing, shaping, or queue management, it can consume available bandwidth and

queue capacity, leaving TCP applications with progressively less effective throughput. TFTP uses UDP port 69 for its initial request exchange and subsequently uses dynamically selected UDP ports for data transfer. SNMP commonly uses UDP port 161 for queries and UDP port 162 for notifications. Although these protocols are often low-volume in normal operations, excessive polling, large transfers, bursts, or poorly managed traffic can produce the UDP behavior that contributes to TCP starvation. Cisco describes the importance of congestion avoidance and differentiated handling for TCP and UDP traffic in its QoS documentation: [Cisco IOS QoS Congestion Avoidance Configuration Guide](#). The transport assignments for these protocols are documented by the IANA: [IANA Service Name and Transport Protocol Port Number Registry](#).

QUESTION NO: 22

Refer to the exhibit.



A network engineer finds that PC1 is accessing the hotel website to do the booking but fails to make payment. Which action resolves the issue?

- A. Allow stub network 10.10.202.168/30 on router R3 OSPF.
- B. Decrease the AD to 5 OSPF route 192.168.94.0 on R1.
- C. Increase the AD to 200 of static route 192.168.94.0 on R3.
- D. Configure a reverse route on R1 for PC1 172.16.1.0/24.

ANSWER: A

Explanation:

Allowing stub network 10.10.202.168/30 on router R3 OSPF resolves the issue because that connected subnet must be included in the OSPF domain so that all participating routers have a consistent, reachable path to the payment-related destination and a valid return path toward PC1. The fact that the booking portion of the hotel site is reachable indicates that basic client connectivity exists, but the payment workflow can use a different server or network path. If the R3-connected /30 network is omitted from OSPF advertisement, routers elsewhere in the topology may lack the route needed to forward traffic through that segment or to return payment-session traffic correctly. Advertising the subnet through OSPF causes R3 to originate the connected prefix in its link-state information, allowing the SPF calculation to install the appropriate routes throughout the area. This restores bidirectional IP reachability required by the payment transaction. Cisco documents that OSPF interfaces and associated networks are enabled through OSPF configuration and advertised as part of the OSPF routing domain; see the [Cisco IOS OSPF Configuration Guide](#). The OSPF protocol behavior for originating and distributing router link-state information is also defined in [RFC 2328](#).

QUESTION NO: 23

```
snmp-server community ciscotest1
snmp-server host 192.168.1.128 ciscotest
snmp-sever enable traps bgp
```

Refer to the exhibit. Network operations cannot read or write any configuration on the device with this configuration from the operations subnet.

Which two configurations fix the issue? (Choose two.)

- A. Configure SNMP rw permission in addition to community ciscotest.
- B. Modify access list 1 and allow operations subnet in the access list.
- C. Modify access list 1 and allow SNMP in the access list.
- D. Configure SNMP rw permission in addition to version 1.
- E. Configure SNMP rw permission in addition to community ciscotest 1.

ANSWER: B E

Explanation:

The configuration must both authorize the operations subnet as an SNMP manager and grant that manager write-level access. Modifying access list 1 to permit the operations subnet resolves the source-address restriction associated with the SNMP community. When an access list number is appended to an SNMP community definition, only SNMP requests sourced from addresses permitted by that list can use the community string. The operations subnet therefore needs a matching permit entry in access list 1.

Configuring SNMP rw permission in addition to community ciscotest 1 establishes the required read-write community access while retaining the access-list association. In IOS syntax, this is represented by an SNMP community configuration equivalent to `snmp-server community ciscotest rw 1`. Read-write permission allows SNMP SET operations, which are necessary for supported configuration changes through SNMP; the same community also supports SNMP reads. Together, the access-list permit and the read-write community setting ensure that requests from the authorized operations subnet can authenticate with the community and perform both monitoring and permitted management writes. Cisco describes community-string access control and read-only/read-write authorization in its [SNMP configuration tutorial](#) and [SNMP community-string documentation](#).

QUESTION NO: 24

What are two characteristics of OSPF not-so-stubby areas? (Choose two.)

- A. An NSSA ASBR generates Type 7 LSAs for redistributed routes.
- B. An NSSA blocks all Type 3 summary LSAs.
- C. An NSSA ABR translates eligible Type 7 LSAs into Type 5 LSAs.
- D. An NSSA requires all routers to be configured as ASBRs.
- E. An NSSA permits Type 5 LSAs to be flooded from other areas.

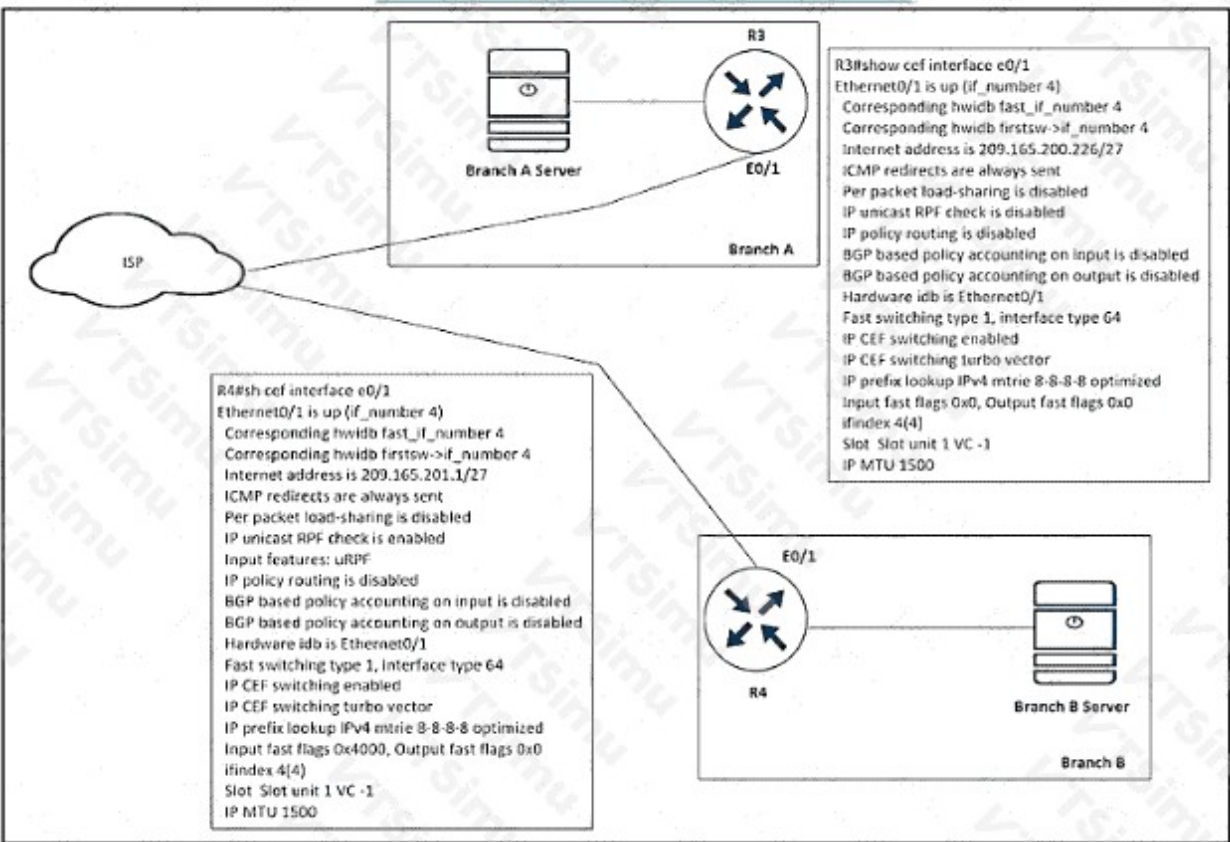
ANSWER: A C

Explanation:

An NSSA permits an ASBR inside the area to redistribute external routes. The ASBR advertises these redistributed prefixes as Type 7 LSAs rather than Type 5 LSAs within the NSSA. Therefore, **an NSSA ASBR generates Type 7 LSAs for redistributed routes** is correct.

The ABR connected to the NSSA translates eligible Type 7 LSAs into Type 5 external LSAs before advertising them into normal OSPF areas. This permits external reachability to be distributed from the NSSA to the rest of the OSPF domain while preserving the NSSA flooding rules. NSSA operation is described in [RFC 3101](#) and Cisco documents configuration details in the [Cisco IOS XE OSPF Configuration Guide](#).

QUESTION NO: 25



Refer to the exhibit.

A shoe retail company implemented the uRPF solution for an antispoofing attack. A network engineer received the call that the branch A server is under an IP spoofing attack. Which configuration must be implemented to resolve the attack?

A)

```
R4
interface ethernet0/1
ip unicast RPF check reachable-via any allow-default allow-self-ping
```

B)

```
R4
interface ethernet0/1
ip verify unicast source reachable-via any allow-default allow-self-ping
```

C)

```
R3
interface ethernet0/1
ip verify unicast source reachable-via any allow-default allow-self-ping
```

D)

```
R3
interface ethernet0/1
 ip unicast RPF check reachable-via any allow-default allow-self-ping
```

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

ANSWER: C

Explanation:

Option C is correct because uRPF must validate packets as they enter the routing device from the direction in which spoofed traffic can arrive before those packets are forwarded toward the Branch A server. The configuration applies reverse-path validation on the applicable ingress interface, causing the router to look up the packet source address in its routing information and verify that the return path is valid for that receiving interface. Packets whose claimed source has no valid reverse route are discarded, which prevents an attacker from successfully using an unauthorized or forged source address to reach the server.

This placement is important: uRPF is an ingress filtering mechanism, so it protects the server by stopping invalid-source traffic at the network edge nearest the incoming attack path rather than after the traffic has already crossed the internal network. Strict reverse-path validation is appropriate where routing to the source is expected to be symmetric; it provides the strongest source-address verification in that situation. Cisco documents uRPF as a method for dropping packets that fail reverse-path source validation, and RFC 3704 describes this reverse-path approach as a core method of ingress filtering against source-address spoofing. See [Cisco: Unicast Reverse Path Forwarding](#) and [RFC 3704: Ingress Filtering for Multihomed Networks](#).

QUESTION NO: 26

What are two functions of IPv6 Source Guard? (Choose two.)

- A. It works independent from IPv6 neighbor discovery.
- B. It denies traffic from unknown sources or unallocated addresses.
- C. It uses the populated binding table to allow legitimate traffic.
- D. It denies traffic by inspecting neighbor discovery packets for specific patterns.
- E. It blocks certain traffic by inspecting DHCP packets for specific sources.

ANSWER: B C

Explanation:

IPv6 Source Guard is a first-hop security feature that prevents IPv6 address spoofing on untrusted Layer 2 access ports. It enforces source-address validation by using an IPv6-to-MAC-address binding associated with the switch interface. The binding information is typically learned from DHCPv6 snooping and can also be created through IPv6 Neighbor Discovery–based binding mechanisms or static configuration, depending on the platform and deployment. Therefore, *It uses the populated binding table to allow legitimate traffic.* describes its core operation: traffic is permitted when its source information matches a valid binding for that port.

Because forwarding is restricted to valid, learned bindings, *It denies traffic from unknown sources or unallocated addresses.* is also correct. A host cannot simply select and transmit using an arbitrary IPv6 source address when IPv6 Source Guard is enabled on the access interface; packets whose source does not match an authorized binding are blocked. This

enforcement limits spoofing and helps ensure that hosts use addresses legitimately assigned or verified for their connected port. Cisco describes IPv6 Source Guard as filtering IPv6 traffic based on the source IPv6 and MAC address binding table. See [Cisco IPv6 Source Guard documentation](#) and [Cisco IPv6 First-Hop Security overview](#).

QUESTION NO: 27

Which two statements describe TACACS+ operation on Cisco IOS XE devices? (Choose two.)

- A. It separates authentication, authorization, and accounting functions.
- B. It uses TCP port 49 for client-server communication.
- C. It encrypts only the password field of each request.
- D. It uses UDP port 1812 for all AAA transactions.
- E. It cannot provide command authorization.

ANSWER: A B

Explanation:

TACACS+ separates AAA functions into authentication, authorization, and accounting. This permits an administrator to authenticate a user, authorize the commands or privilege level available to that user, and record management activity independently. Therefore, **it separates authentication, authorization, and accounting functions** is correct.

TACACS+ uses TCP port 49 for client-server communication. The reliable TCP transport supports the exchange of AAA requests and responses between the network device and the TACACS+ server. RADIUS is a different AAA protocol and commonly uses UDP rather than TCP. Cisco provides TACACS+ configuration guidance in the [Cisco IOS XE AAA Configuration Guide](#).

QUESTION NO: 28

What is the output of the following command:

```
show ip vrf
```

- A. Displays VRF names, default RD values, and associated interfaces.
- B. Displays IP routing table information associated with a VRF
- C. Show 's routing protocol information associated with a VRF.
- D. Displays the ARP table (static and dynamic entries) in the specified VRF

ANSWER: A

Explanation:

The `show ip vrf` command displays the configured IPv4 VRFs on the device. Its standard output includes the VRF name, the default route distinguisher (RD) associated with each VRF, and the interfaces assigned to that VRF. The RD is a value used in MPLS Layer 3 VPN deployments to make otherwise overlapping IPv4 prefixes unique when they are represented as VPNv4 routes. Even where an RD is not explicitly configured, the command presents the VRF's default RD field as part of its summary output. This makes the command useful for quickly verifying that VRFs exist, confirming their RD values, and checking interface-to-VRF associations before validating more detailed routing or forwarding behavior. Cisco's VRF documentation describes VRFs as separate routing and forwarding instances, each capable of maintaining distinct routing information, and Cisco command references identify `show ip vrf` as the operational command for displaying VRF information. See the [Cisco VRF configuration guide](#) and the [Cisco IP Routing VRF command reference](#).

QUESTION NO: 29

A network engineer must configure a DMVPN network so that a spoke establishes a direct path to another spoke if the two must send traffic to each other. A spoke must send traffic directly to the hub if required. Which configuration meets this requirement?

A)

```
At the hub router
interface tunnel10
ip nhrp nhs multicast dynamic
ip nhrp redirect
tunnel mode gre multipoint

On the spokes router
interface tunnel10
ip nhrp nhs multicast dynamic
ip nhrp shortcut
tunnel mode gre multipoint
```

B)

```
At the hub router
interface tunnel10
ip nhrp nhs dynamic multipoint
ip nhrp nhs shortcut
tunnel mode gre multicast

On the spokes router
interface tunnel10
ip nhrp nhs multicast dynamic
ip nhrp nhs redirect
tunnel mode gre multipoint
```

C)

```
At the hub router
interface tunnel10
ip nhrp nhs multicast dynamic
ip nhrp nhs shortcut
tunnel mode gre multipoint

On the spokes router
interface tunnel10
ip nhrp nhs multicast dynamic
ip nhrp nhs redirect
tunnel mode gre multipoint
```

D)

```
At the hub router
interface tunnel10
ip nhrp nhs multicast multipoint
ip nhrp redirect
tunnel mode gre multicast
session (None) (greobnct)
```

A. Option

B. Option

C. Option

D. Option

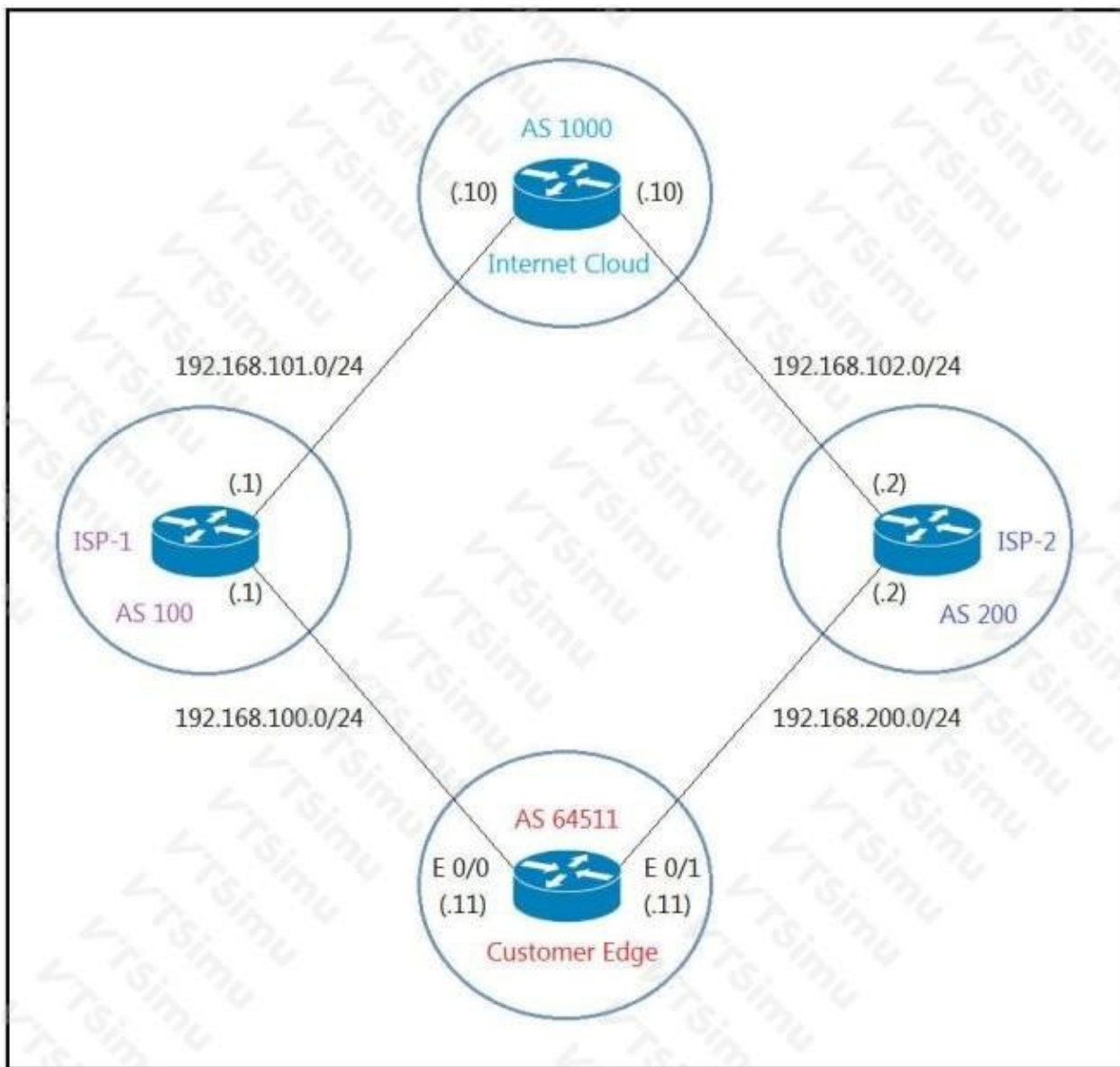
ANSWER: A

Explanation:

The configuration shown in IMAGE_1 meets the requirement because it implements the DMVPN behavior needed for dynamic spoke-to-spoke forwarding while retaining direct spoke-to-hub connectivity. In a DMVPN Phase 3 design, the hub uses NHRP redirect capability and spokes use NHRP shortcut capability. When a spoke initially sends traffic toward another spoke, the packet can be forwarded through the hub according to the routing table. The hub then provides NHRP information that enables the originating spoke to resolve the destination spoke's NBMA address and build a dynamic GRE/IPsec tunnel directly to it. Subsequent packets use that direct spoke-to-spoke tunnel rather than unnecessarily traversing the hub.

The spoke still has an NHRP next-hop server relationship with the hub and a static NHRP mapping for the hub's NBMA address. Therefore, traffic whose destination is the hub itself, or traffic that must use the hub as its next hop, is sent directly to the hub across the established DMVPN tunnel. This combination provides scalable on-demand spoke tunnels without losing the hub's control-plane and forwarding role. Cisco describes NHRP redirects and shortcuts as the mechanisms used to optimize spoke-to-spoke traffic in DMVPN Phase 3: [Cisco DMVPN Overview](#).

QUESTION NO: 30



Refer to the exhibit. The network administrator has configured the Customer Edge router (AS 64511) to send only summarized routes toward ISP-1 (AS 100) and ISP-2 (AS 200).

```
router bgp 64511
network 172.16.20.0 mask 255.255.255.0
network 172.16.21.0 mask 255.255.255.0
network 172.16.22.0 mask 255.255.255.0
network 172.16.23.0 mask 255.255.255.0
aggregate-address 172.16.20.0 255.255.252.0
```

After this configuration, ISP-1 and ISP-2 continue to receive the specific routes and the summary route.

Which configuration resolves the issue?

- A. **router bgp 64511**
aggregate-address 172.16.20.0 255.255.252.0 summary-only
- B. **router bgp 64511**
neighbor 192.168.100.1 summary-only
neighbor 192.168.200.2 summary-only
- C. **ip prefix-list PL_BLOCK_SPECIFIC deny 172.16.20.0/22 ge 22**
ip prefix-list PL_BLOCK_SPECIFIC permit 172.16.20.0/22
!
route-map BLOCK_SPECIFIC permit 10
match ip address prefix-list PL_BLOCK_SPECIFIC
!
router bgp 64511
aggregate-address 172.16.20.0 255.255.252.0 suppress-map BLOCK_SPECIFIC

- D. **interface E 0/0**
ip bgp suppress-map BLOCK_SPECIFIC
!
interface E 0/1
ip bgp suppress-map BLOCK_SPECIFIC
!
ip prefix-list PL_BLOCK_SPECIFIC permit 172.16.20.0/22 ge 24
!
route-map BLOCK_SPECIFIC permit 10
match ip address prefix-list PL_BLOCK_SPECIFIC

A. Option A

- B. Option B
- C. Option C
- D. Option D
- E. Configure `aggregate-address summary-only` under `address-family ipv4 unicast` for BGP AS 64511.

ANSWER: E

Explanation:

The required configuration is an IPv4 BGP aggregate route configured with the `summary-only` keyword in the address family that advertises routes to the ISPs. Cisco BGP normally advertises both an aggregate and its more-specific component prefixes when an aggregate is created. Applying `summary-only` causes BGP to suppress advertisement of the component routes while continuing to originate and advertise the aggregate. Thus, both ISP-1 and ISP-2 receive only the intended summarized prefix, provided that the aggregate has contributing routes in the BGP table.

The aggregate statement must use the actual aggregate network and mask from the exhibit. It must also be placed under `address-family ipv4 unicast` on platforms using address-family BGP configuration. For example, the relevant form is `aggregate-address <summary-network> <summary-mask> summary-only`. This behavior is applied to BGP route advertisement and does not require separate outbound filters merely to hide the component routes. Cisco documents the `summary-only` aggregate-address behavior in its [BGP command reference](#). Additional aggregation concepts are described in the [Cisco BGP aggregation guide](#).

QUESTION NO: 31

Which two statements describe EIGRP feasible successors? (Choose two.)

- A. It is maintained in the EIGRP topology table as a backup path.
- B. It must have a reported distance lower than the feasible distance.
- C. It is installed in the routing table with the same administrative distance as the successor.
- D. It is selected only when EIGRP is configured for unequal-cost load balancing.
- E. It must have a metric that is identical to the successor metric.

ANSWER: A B

Explanation:

A feasible successor is a loop-free backup route that is retained in the EIGRP topology table. It is not installed in the routing table while the successor route remains available. Therefore, **it is maintained in the EIGRP topology table as a backup path** is correct.

For a neighbor route to qualify as a feasible successor, the neighbor-reported distance must be lower than the local router's feasible distance for the destination. This feasibility condition ensures that the neighbor is not routing through the local router to reach that network and provides loop-free backup forwarding. Cisco describes EIGRP route selection and the feasibility condition in the [Cisco IOS XE EIGRP Configuration Guide](#).

QUESTION NO: 32

Which list defines the contents of an MPLS label?

- A. 20-bit label; 3-bit traffic class; 1-bit bottom stack; 8-bit TTL
- B. 32-bit label; 3-bit traffic class; 1-bit bottom stack; 8-bit TTL

C. 20-bit label; 3-bit flow label; 1-bit bottom stack; 8-bit hop limit

D. 32-bit label; 3-bit flow label; 1-bit bottom stack; 8-bit hop limit

ANSWER: A

Explanation:

The MPLS label stack entry is a fixed 32-bit field made up of four subfields: a 20-bit Label value, a 3-bit Traffic Class field, a 1-bit Bottom of Stack field, and an 8-bit Time to Live field. Therefore, *20-bit label; 3-bit traffic class; 1-bit bottom stack; 8-bit TTL* correctly defines the contents and sizes of the fields in one MPLS label stack entry.

The Label field identifies the forwarding treatment associated with the packet at an MPLS-capable router. The Traffic Class field, previously called EXP, carries traffic-class and QoS-related information. The Bottom of Stack bit indicates whether this entry is the final label in the stack; when set, the payload following it is no longer another MPLS label stack entry. The TTL field supports loop prevention in MPLS forwarding, similarly to TTL processing in IP. Collectively, these fields total 32 bits, but the label value itself is specifically 20 bits—not 32 bits. RFC 5462 formally documents the Traffic Class terminology while retaining this label stack entry format. See [RFC 5462](#) and the base MPLS label-stack encoding in [RFC 3032](#).

QUESTION NO: 33

Refer to the exhibit.

```
RR# show running-config
!
interface Ethernet0/1
  no ip address
  ipv6 address 2001:DB8:1:12::2/64
  ipv6 traffic-filter ACL in
!
ipv6 access-list ACL
sequence 10 permit tcp any any eq 22
sequence 20 permit tcp any eq 22 any
sequence 30 permit tcp any any eq bgp
sequence 40 permit tcp any eq bgp any
sequence 50 permit udp any any eq ntp
sequence 60 permit udp any eq ntp any
sequence 70 permit udp any any eq snmp
sequence 80 deny ipv6 any any log

RR# show ipv6 cef ::/0
::/0
  nexthop 2001:DB8:1:12::1 Ethernet0/1

*Feb 23 00:23:17.211: %IPV6_ACL-6-ACCESSLOGDP: list ACL/80
denied icmpv6 2001:DB8:1:12::1 -> FF02::1:FF00:2 (135/0), 7321
packets
```

After a security audit, the administrator implemented an ACL in the route reflector. The RR became unreachable from any router in the network. Which two actions resolve the issue? (Choose two.)

A. Enable the ND proxy feature on the default gateway.

B. Configure a link-local address on the Ethernet0/1 interface.

C. Permit ICMPv6 neighbor discovery traffic in the ACL.

D. Remove the ACL entry 80.

E. Change the next hop of the default route to the link-local address of the default gateway.

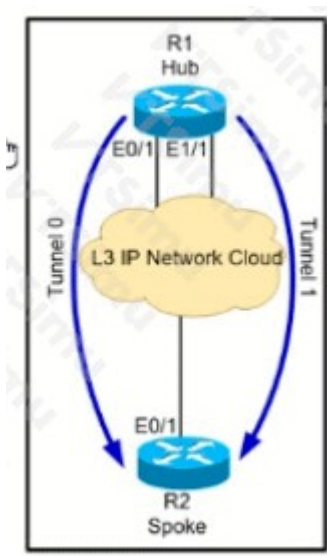
ANSWER: C D

Explanation:

IPv6 Neighbor Discovery is required for devices on an Ethernet segment to resolve a neighbor's Layer 2 address and maintain neighbor reachability. It uses ICMPv6 Neighbor Solicitation and Neighbor Advertisement messages, including the ICMPv6 types used for address resolution and neighbor-cache maintenance. Therefore, **Permit ICMPv6 neighbor discovery traffic in the ACL**, restores the control-plane traffic needed for routers to forward packets to the route reflector and receive return traffic from it.

Remove the ACL entry 80. is also valid because an explicit broad deny statement placed before the implicit end-of-ACL behavior blocks Neighbor Discovery messages. Cisco IPv6 ACL processing includes implicit permits for Neighbor Discovery Neighbor Solicitation and Neighbor Advertisement messages only when those packets are not already matched and denied by an explicit ACE. Removing the explicit deny allows the standard implicit Neighbor Discovery handling to take effect while preserving the intended filtering entries above it. Neighbor Discovery's function and required message types are defined in [RFC 4861](#). Cisco documents the implicit Neighbor Discovery handling and IPv6 traffic-filter behavior in its [IPv6 Traffic Filtering Configuration Guide](#).

QUESTION NO: 34



Refer to the exhibit. The hub and spoke are connected via two DMVPN tunnel interfaces. The NHRP is configured and the tunnels are detected on the hub and the spoke. Which configuration command adds an IPsec profile on both tunnel interfaces to encrypt traffic?

- A. tunnel protection ipsec profile DMVPN multipoint
- B. tunnel protection ipsec profile DMVPN tunnel1
- C. tunnel protection ipsec profile DMVPN shared
- D. tunnel protection ipsec profile DMVPN unique

ANSWER: C

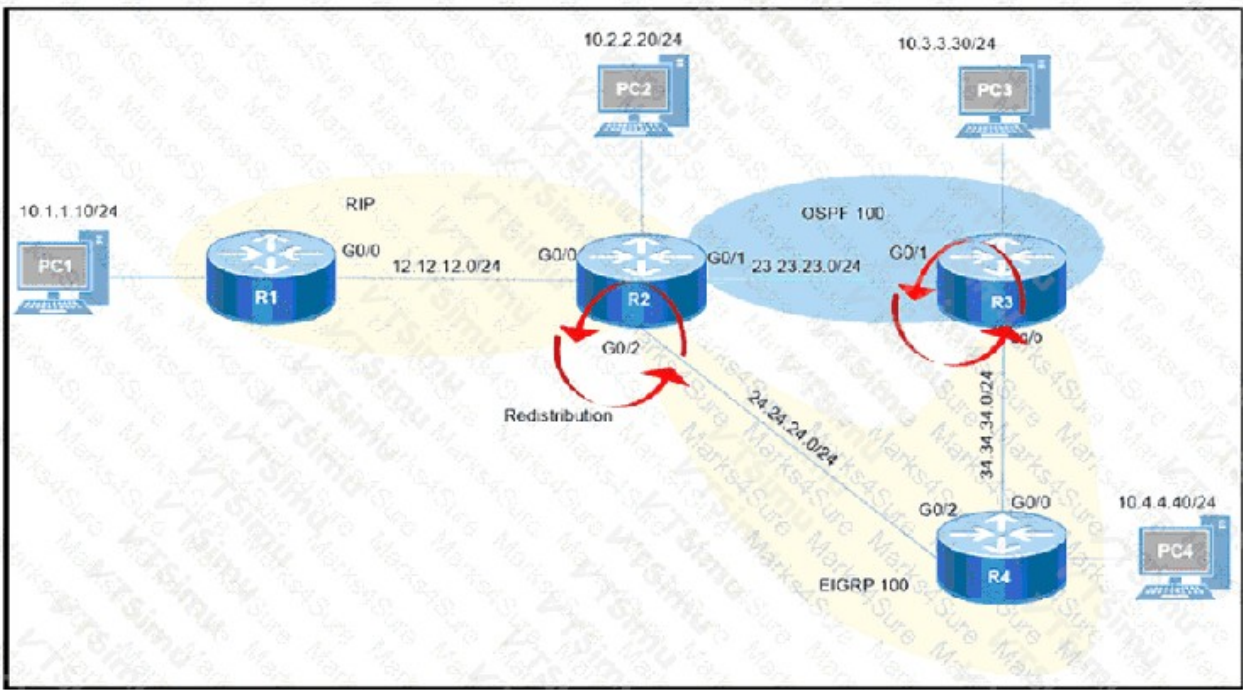
Explanation:

`tunnel protection ipsec profile DMVPN shared` is the correct tunnel-interface command when the same IPsec profile must protect more than one DMVPN tunnel interface. The profile name is `DMVPN`, which identifies the previously configured IPsec profile containing the required transform-set or IPsec-proposal settings. The `shared` keyword is significant

in a dual-tunnel DMVPN design: it permits that IPsec profile to be shared across tunnel interfaces rather than associating its security-association handling with only one interface. Applying this command under each DMVPN tunnel interface enables IPsec protection for GRE/DMVPN traffic while allowing NHRP-based dynamic spoke-to-spoke operation to continue. This is configured in tunnel-interface configuration mode on both the hub and spoke wherever the corresponding DMVPN tunnels terminate. Cisco documents the `tunnel protection ipsec profile interface` command and its `shared` keyword in its IOS command reference, and its DMVPN configuration guidance describes using tunnel protection to secure DMVPN overlays. See the [Cisco DMVPN tunnel protection documentation](#) and the [Cisco IOS DMVPN configuration guide](#).

QUESTION NO: 35

Refer to the exhibit.



Redistribution is enabled between the routing protocols, and now PC2, PC3, and PC4 cannot reach PC1. What are the two solutions to fix the problem? (Choose two.)

- A. Filter RIP routes back into RIP when redistributing into RIP in R2
- B. Filter OSPF routes into RIP FROM EIGRP when redistributing into RIP in R2.
- C. Filter all routes except RIP routes when redistributing into EIGRP in R2.
- D. Filter RIP AND OSPF routes back into OSPF from EIGRP when redistributing into OSPF in R2
- E. Filter all routes except EIGRP routes when redistributing into OSPF in R3.

ANSWER: A C

Explanation:

Mutual redistribution must prevent a route that originated in one routing domain from returning to that same domain through another redistribution path. In this topology, the route for PC1 originates in RIP. Without filtering, that RIP-originated prefix can be redistributed into EIGRP or OSPF and then be redistributed back into RIP. The returned route can be selected or propagated as a redistributed copy rather than preserving the intended original path, creating a redistribution feedback loop and causing loss of reachability.

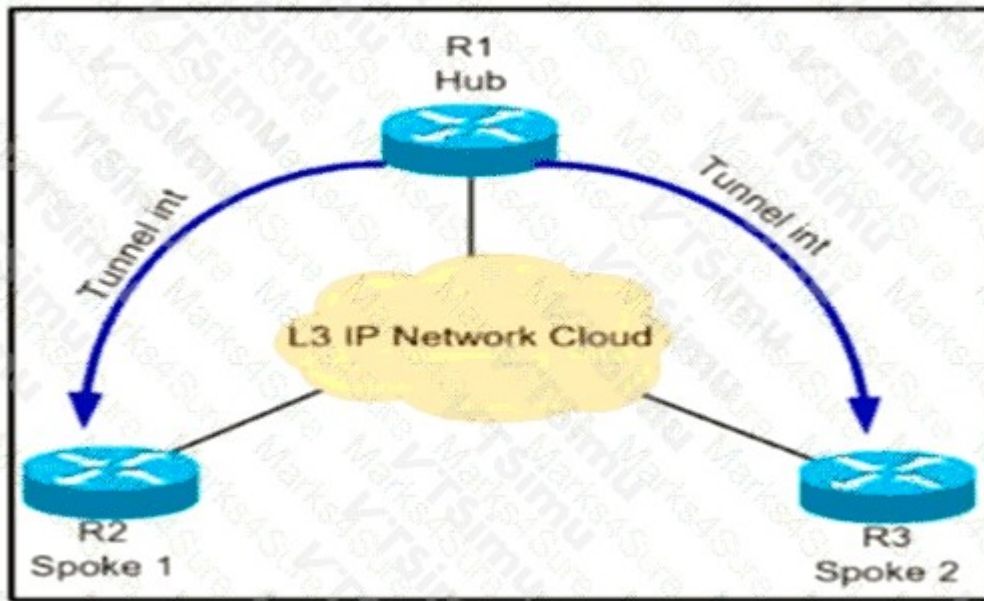
Filtering RIP routes back into RIP when redistributing into RIP in R2 prevents RIP-originated prefixes that have traversed another protocol from being reintroduced into their source protocol. Filtering all routes except RIP routes when redistributing into EIGRP in R2 ensures that EIGRP receives only the RIP-originated routes intended for that redistribution direction; it blocks routes that originated in EIGRP or arrived from another redistributed domain from being fed back into EIGRP. In

production, Cisco recommends using route tags with route maps to identify redistributed routes and deny them on a return redistribution path, rather than matching only prefixes. See Cisco's [route redistribution guidance](#) and [RIP redistribution documentation](#).

QUESTION NO: 36

Refer to Exhibit.

A network administrator has successfully configured DMVPN topology between a hub and two spoke routers. Which two configuration commands should establish direct communications between spoke 1 and spoke 2 without going through the hub? (Choose two).



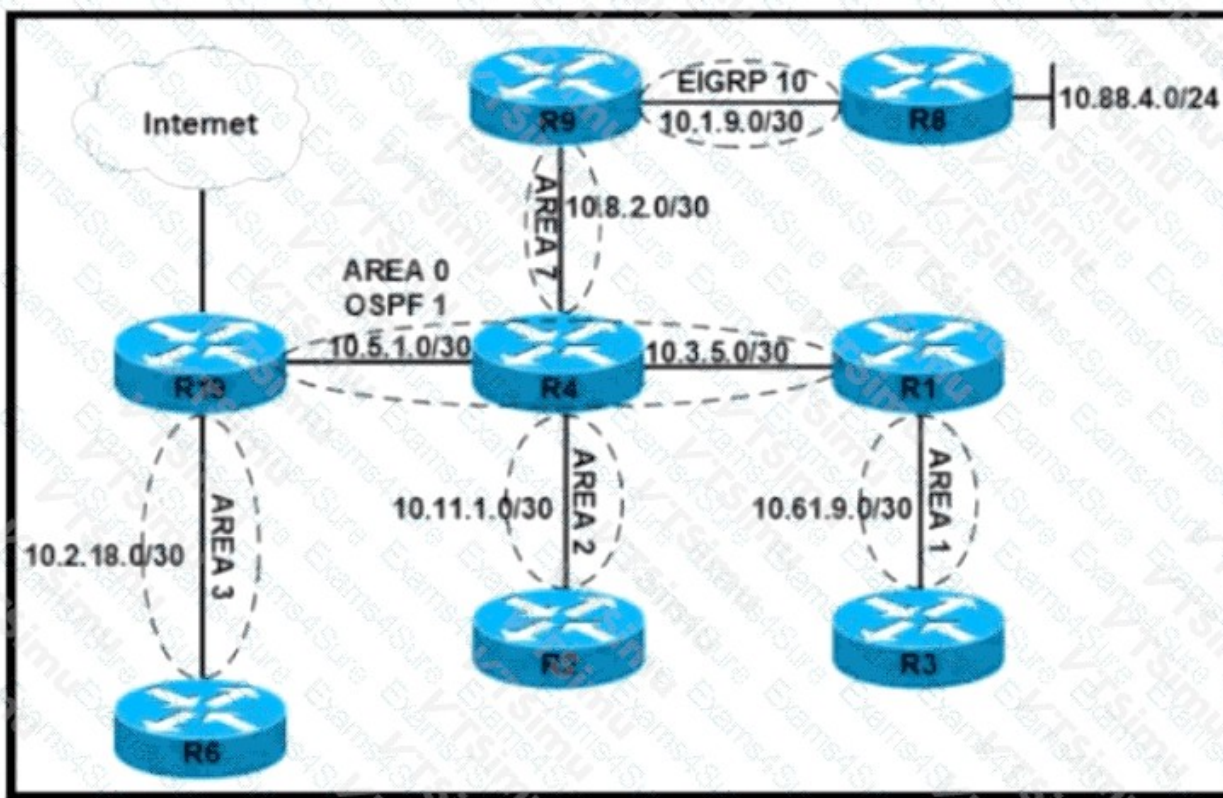
- A. At the hub router, configure the `ip nhrp shortcut` command.
- B. At the spoke routers, configure the `ip nhrp spoke-tunnel` command.
- C. At the hub router, configure `ip nhrp redirect` the command
- D. At the spoke routers, configure the `ip nhrp shortcut` command.
- E. At the hub router, configure the `ip nhrp spoke-tunnel` command

ANSWER: C D

Explanation:

Direct dynamic spoke-to-spoke forwarding in a DMVPN deployment is provided by DMVPN Phase 3 using NHRP redirect and shortcut functionality. The hub tunnel interface must use `ip nhrp redirect`. When a packet from one spoke initially reaches the hub while destined for a network behind another spoke, the hub sends an NHRP redirect to the originating spoke. This informs that spoke that a more optimal path is available and triggers resolution of the destination spoke's tunnel mapping.

Each spoke tunnel interface must use `ip nhrp shortcut`. This permits the spoke to process the redirect and install a CEF shortcut entry that replaces the hub as the data-plane next hop with the destination spoke's tunnel address. The spoke then dynamically builds an on-demand GRE/IPsec tunnel directly to the other spoke using the resolved NBMA address. Subsequent traffic uses that direct tunnel rather than being forwarded through the hub. These commands are complementary: redirect is the hub's notification mechanism, while shortcut enables the spokes' route and forwarding-table optimization. Cisco documents these DMVPN Phase 3 NHRP functions in its DMVPN configuration guidance: [Cisco IOS DMVPN Configuration Guide](#) and [Cisco NHRP and DMVPN documentation](#).



Refer to the exhibit After an engineer modified the configuration for area 7 to permit type 1, 2, and 7 LSAs only users connected to router R9 reported that they could no longer access the internet. Which configuration restores internet access to users on R9 and permits only LSA type 1, 2, and 7?

A)

```
R4#
router ospf 1
 area 0 nssa default-information-originate
 network 10.5.1.0 0.0.0.3 area 0
 network 10.8.2.0 0.0.0.3 area 7

R9#
router ospf 1
 area 7 nssa
 redistribute eigrp 10 subnets
 network 10.8.2.0 0.0.0.3 area 7
```

B)

```
R4#
router ospf 1
  area 7 nssa no-summary
  network 10.5.1.0 0.0.0.3 area 0
  network 10.8.2.0 0.0.0.3 area 7
```

```
R9#
router ospf 1
  area 7 nssa
  redistribute eigrp 10 subnets
  network 10.8.2.0 0.0.0.3 area 7
```

C)

```
R4#
router ospf 1
  area 7 nssa
  network 10.5.1.0 0.0.0.3 area 0
  network 10.8.2.0 0.0.0.3 area 7
```

```
R9#
router ospf 1
  area 7 nssa
  redistribute eigrp 10 subnets
  network 10.8.2.0 0.0.0.3 area 7
```

D)

```
R4#
router ospf 1
  area 0 area 7 stub no-summary
  network 10.5.1.0 0.0.0.3 area 0
  network 10.8.2.0 0.0.0.3 area 7
```

```
R9#
router ospf 1
  area 7 stub
  redistribute eigrp 10 subnets
  network 10.8.2.0 0.0.0.3 area 7
```

- A. Option A
- B. area 7 nssa no-summary
- C. Option C
- D. Option D

ANSWER: B

Explanation:

The `area 7 nssa no-summary` configuration creates a totally NSSA. Within that area, routers retain intra-area topology information as type 1 and type 2 LSAs, and external information is represented with type 7 LSAs. The `no-summary` keyword

suppresses interarea summary advertisements, ensuring that type 3 and type 4 LSAs are not flooded into area 7. This meets the stated LSA restriction while retaining the NSSA behavior required for external-route handling.

Importantly, an ABR configured for a totally NSSA originates a default route into the NSSA as a type 7 LSA. R9 can therefore install and use that default route to forward Internet-bound traffic to the ABR, restoring connectivity without permitting the unwanted LSA types. The type 7 default is local to the NSSA; when appropriate, NSSA external routes can be translated by the ABR for advertisement toward other OSPF areas. Cisco documents that an NSSA ABR using `no-summary` injects a type 7 default route, and that NSSA external information is carried in type 7 LSAs. See Cisco's [OSPF LSA documentation](#) and [NSSA configuration guide](#).

QUESTION NO: 38

Refer to the exhibit.

```
TAC+: TCP/IP open to 171.68.118.101/49 failed --  
Destination unreachable; gateway or host down  
AAA/AUTHEN (2546660185): status = ERROR  
AAA/AUTHEN/START (2546660185): Method=LOCAL  
AAA/AUTHEN (2546660185): status = FAIL  
As1 CHAP: Unable to validate Response. Username chapuser: Authentication failure
```

Why is user authentication being rejected?

- A. The TACACS+ server expects "user", but the NT client sends "domain/user".
- B. The TACACS+ server refuses the user because the user is set up for CHAP.
- C. The TACACS+ server is down, and the user is in the local database.
- D. The TACACS+ server is down, and the user is not in the local database.

ANSWER: D

Explanation:

The TACACS+ server is down, and the user is not in the local database. The debug output indicates that the device cannot obtain a usable TACACS+ authentication response, which causes the configured AAA method list to proceed to its local authentication fallback. AAA fallback is evaluated in method-list order: when the TACACS+ service is unavailable rather than explicitly rejecting credentials, IOS can try the next method, commonly `local`. Local authentication requires a matching locally configured username and the appropriate password or secret. Because no matching local user exists, IOS has no successful authentication source available and rejects the login attempt. This behavior is important operationally: a local break-glass account can preserve administrative access during a TACACS+ outage, but it must be explicitly configured and included after the TACACS+ group in the relevant AAA method list. Cisco's TACACS+ PPP troubleshooting guidance discusses using debug output to distinguish server communication failures from authentication exchanges, while the TACACS+ protocol specification defines the authentication communication model between a network access device and server. See [Cisco TACACS+ PPP troubleshooting](#) and [RFC 8907: The TACACS+ Protocol](#).

QUESTION NO: 39

An engineer needs dynamic routing between two routers and is unable to establish OSPF adjacency. The output of the `show ip ospf neighbor` command shows that the neighbor state is EXSTART/ EXCHANGE.

Which action should be taken to resolve this issue?

- A. match the passwords
- B. match the hello timers

- C. match the MTUs
- D. match the network types

ANSWER: C

Explanation:

OSPF neighbors in the EXSTART/EXCHANGE state have successfully progressed through initial neighbor discovery and are trying to negotiate database synchronization. At this stage, the routers exchange Database Description (DBD) packets, which describe the link-state advertisements each router has in its LSDB. By default, OSPF checks the interface MTU advertised in received DBD packets against the local interface MTU. If the values differ, the router can reject the DBD packet and the adjacency commonly remains stuck in EXSTART or EXCHANGE rather than reaching FULL. Therefore, the appropriate corrective action is to match the MTUs on the interfaces connecting the two OSPF routers. Verify the effective Layer 3 interface MTU on both ends and correct any mismatch, including one caused by a tunnel, subinterface, or provider handoff. Cisco documents MTU mismatch as a primary cause of OSPF neighbors stuck in EXSTART/EXCHANGE and recommends making the MTU values consistent. Although the `ip ospf mtu-ignore` interface command can bypass the OSPF MTU check in limited cases, matching the actual interface MTUs is the preferred operational fix because it avoids packets exceeding the smaller path MTU during LSDB synchronization or later traffic forwarding. See Cisco's [OSPF neighbor troubleshooting guide](#) and the [Cisco IOS OSPF command reference](#).

QUESTION NO: 40

Which protocol does MPLS use to support traffic engineering?

- A. Tag Distribution Protocol (TDP)
- B. Resource Reservation Protocol (RSVP)
- C. Border Gateway Protocol (BGP)
- D. Label Distribution Protocol (LDP)

ANSWER: B

Explanation:

Resource Reservation Protocol (RSVP) is correct because MPLS Traffic Engineering uses RSVP with Traffic Engineering extensions, commonly called RSVP-TE, as its signaling protocol. RSVP-TE establishes and maintains traffic-engineered label-switched paths (TE tunnels) across an MPLS network. When signaling a tunnel, it carries TE-specific information such as requested bandwidth, explicit route objects, setup and holding priorities, and affinity constraints. Each transit router evaluates available resources and, when the path can be admitted, reserves the requested bandwidth and assigns the MPLS labels needed to forward traffic along that engineered path. This enables an operator to steer selected traffic away from the normal shortest IGP path and onto a path that meets defined capacity or policy requirements. In Cisco IOS XE implementations, RSVP is enabled on participating interfaces and RSVP-TE performs the signaling required for MPLS TE tunnels. Cisco documents RSVP as the resource-reservation and signaling mechanism used by MPLS TE: [Cisco IOS XE MPLS Traffic Engineering Configuration Guide](#). The RSVP protocol and its resource-reservation model are defined by the IETF in [RFC 2205](#).

QUESTION NO: 41

What are two characteristics of VRF instance? (Choose two.)

- A. All VRFs share customers routing and CEF tables .
- B. An interface must be associated to one VRF.
- C. Each VRF has a different set of routing and CEF tables
- D. It is defined by the VPN membership of a customer site attached to a P device.

E. A customer site can be associated to different VRFs

ANSWER: B C

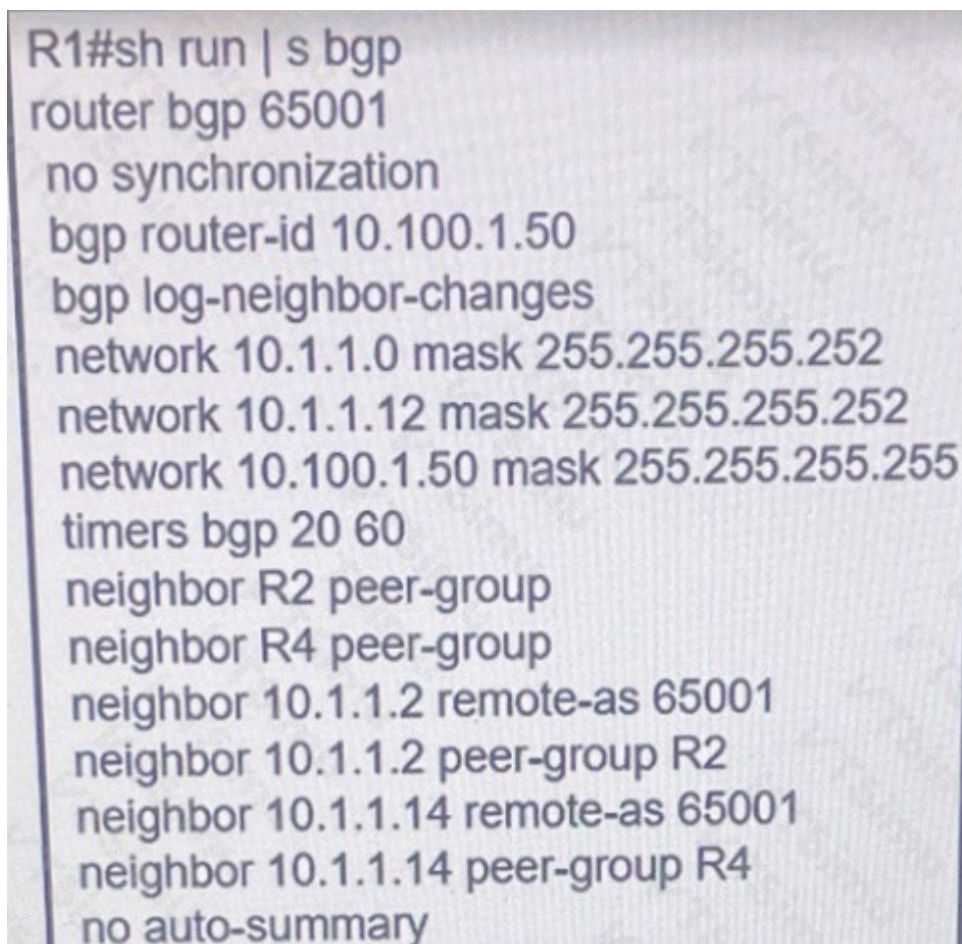
Explanation:

An interface must be associated to one VRF because VRF forwarding separates packet forwarding contexts on a router. When an interface is configured under a VRF forwarding instance, its connected routes and traffic lookup use that instance's forwarding context; an individual Layer 3 interface cannot simultaneously belong to multiple VRF instances. Interfaces that are not assigned to a VRF remain in the global routing table, which is itself separate from configured VRFs.

Each VRF has a different set of routing and CEF tables because a VRF is a separate virtual routing and forwarding domain. It maintains an independent routing-information base and corresponding Cisco Express Forwarding information, allowing overlapping customer IPv4 or IPv6 address space to be used without ambiguity. Routes learned through interfaces in one VRF are installed and resolved within that VRF unless deliberate route leaking or inter-VRF connectivity is configured. This separation is the core mechanism used by MPLS Layer 3 VPNs and VRF-Lite deployments to isolate tenant or departmental traffic on shared network infrastructure. Cisco describes VRF as maintaining separate routing and forwarding tables, and RFC 4364 defines the VPN routing separation model used by provider networks. See [Cisco IOS XE VRF Configuration Guide](#) and [RFC 4364: BGP/MPLS IP Virtual Private Networks](#).

QUESTION NO: 42

Refer to the exhibit.



```
R1#sh run | s bgp
router bgp 65001
no synchronization
bgp router-id 10.100.1.50
bgp log-neighbor-changes
network 10.1.1.0 mask 255.255.255.252
network 10.1.1.12 mask 255.255.255.252
network 10.100.1.50 mask 255.255.255.255
timers bgp 20 60
neighbor R2 peer-group
neighbor R4 peer-group
neighbor 10.1.1.2 remote-as 65001
neighbor 10.1.1.2 peer-group R2
neighbor 10.1.1.14 remote-as 65001
neighbor 10.1.1.14 peer-group R4
no auto-summary
```

While troubleshooting a BGP route reflector configuration, an engineer notices that reflected routes are missing from neighboring routers. Which two BGP configurations are needed to resolve the issue? (Choose two)

- A. neighbor 10.1.1.14 route-reflector-client
- B. neighbor R2 route-reflector-client

- C. neighbor 10.1.1.2 allowas-in
- D. neighbor R4 route-reflector-client
- E. neighbor 10.1.1.2 route-reflector-client

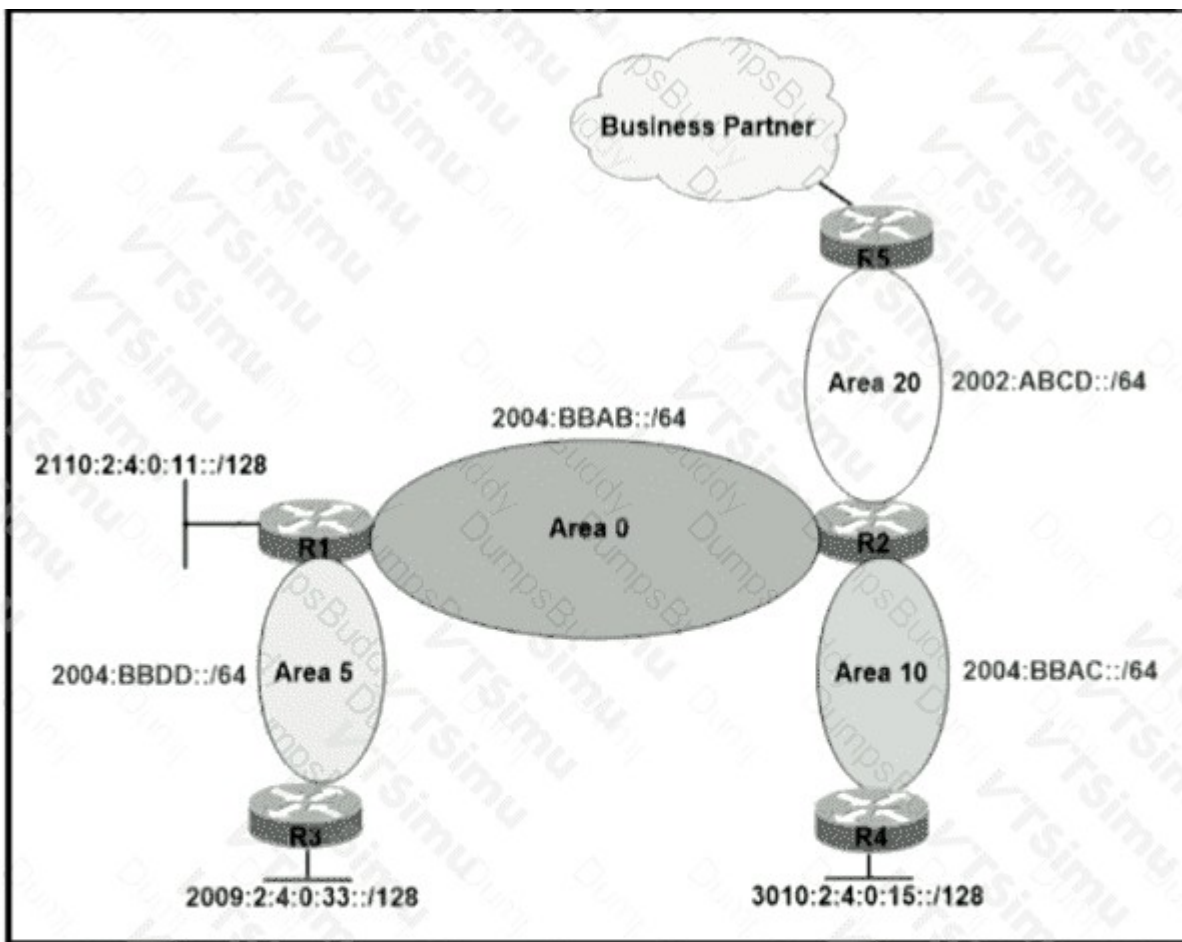
ANSWER: A E

Explanation:

A route reflector must explicitly identify each iBGP peer that is intended to be a route-reflector client. The commands `neighbor 10.1.1.14 route-reflector-client` and `neighbor 10.1.1.2 route-reflector-client` perform that function on the route reflector for the two client peer sessions shown. Once the peers are designated as clients, the reflector can advertise routes learned from one client to the other client and can also reflect routes learned from nonclient peers to its clients. This removes the normal iBGP split-horizon limitation that prevents an iBGP speaker from advertising routes learned through iBGP to another iBGP peer. The commands must use the configured BGP neighbor addresses, because the route-reflector-client setting is applied on a per-neighbor basis within the relevant BGP address-family configuration. Route reflection therefore permits the client routers to receive the missing reflected prefixes without requiring a full mesh of direct iBGP sessions between all routers. Cisco documents route reflection and the client-neighbor configuration model in its [BGP Route Reflector Configuration Guide](#) and describes the per-neighbor [BGP command reference](#).

QUESTION NO: 43

Refer to the exhibit.



```
R2#sh ipv6 route ospf
O 2002:ABCD::/64 [110/1]
  via FastEthernet0/1, directly connected
O 2004:BBAB::/64 [110/1]
  via FastEthernet0/0, directly connected
O 2004:BBAC::/64 [110/1]
  via FastEthernet1/0, directly connected
O 3010:2:4:0:15::/128 [110/1]
  via FE80::C804:1DFF:FE20:8, FastEthernet0/0
```

A network engineer applied a filter for LSA traffic on OSPFv3 interarea routes on the area 5 ABR to protect advertising the internal routes of area 5 to the business partner network. All other areas should receive the area 5 internal routes. After the respective route filtering configuration is applied on the ABR, area 5 routes are not visible on any of the areas. How must the filter list be applied on the ABR to resolve this issue?

- A. in the "in" direction for area 5 on router R1
- B. in the "out" direction for area 5 on router R1
- C. in the "in" direction for area 20 on router R2
- D. in the "out" direction for area 20 on router R2

ANSWER: C

Explanation:

The filter list must be applied **in the "in" direction for area 20 on router R2**. Area 5 prefixes are interarea routes when they reach the ABR connected to the business-partner area. The required policy is not to prevent area 5 routes from leaving their source area; it is to prevent those routes from being injected into the specific partner-facing area while retaining normal advertisement to every other OSPF area. Applying the prefix filter to area 20 in the inbound-to-area direction on R2 scopes

the suppression to the summary route advertisements that would otherwise be sent into area 20. Thus, R2 can continue to advertise the area 5 prefixes through the backbone and toward other eligible areas, but the partner routers in area 20 do not learn those prefixes. OSPF ABRs advertise interarea reachability using summary LSAs, and the area filter-list mechanism is designed to control those interarea advertisements at the area boundary. Cisco documents the OSPFv3 area filter-list configuration and its directional behavior in the [Cisco IOS XE OSPFv3 Configuration Guide](#). The interarea summary-LSA function is also defined in [RFC 5340](#).

QUESTION NO: 44

Refer to Exhibit.

```
ip dhcp excluded-address 172.16.16.1 172.16.16.2
!
ip dhcp pool 0
network 172.16.16.0 255.255.255.0
domain-name cisco.com
dns-server 172.16.16.2
lease 30

interface Ethernet0/0
ip address 10.1.1.1 255.255.255.252
ip access-group 100 in

access-list 100 deny  udp any any
access-list 100 permit ip any any
```

Which two configurations allow clients to get dynamic ip addresses assigned?

- A. Configure access-list 100 permit udp any any eq 61 as the first line
- B. Configure access-list 100 permit udp any any eq 86 as the first line
- C. Configure access-list 100 permit udp any any eq 68 as the first line
- D. Configure access-list 100 permit udp any any eq 69 as the first line
- E. Configure access-list 100 permit udp any any eq 67 as the first line

ANSWER: C E

Explanation:

DHCP address assignment requires bidirectional UDP communication between the client and the DHCP server or relay. A DHCP client initially sends its DHCPDISCOVER and DHCPREQUEST messages using UDP source port 68 and UDP destination port 67. Therefore, placing *Configure access-list 100 permit udp any any eq 67 as the first line* first permits DHCP client requests toward the server-side DHCP service. The DHCP server returns DHCPOFFER and DHCPACK messages with UDP source port 67 and UDP destination port 68. Placing *Configure access-list 100 permit udp any any eq 68 as the first line* first permits these server replies to reach clients, including when the client has not yet received an IP address and relies on broadcast-based DHCP processing. Together, the two destination-port entries permit both directions required for the standard DHCP exchange, allowing a client to complete lease negotiation and receive a dynamically assigned IPv4

address. Cisco documentation identifies UDP ports 67 and 68 as the DHCP server and client ports, respectively; the IETF DHCP specification defines the same client/server port usage. See [Cisco IOS DHCP documentation](#) and [RFC 2131](#).

QUESTION NO: 45

Which two actions are required to implement static route failover by using an IP SLA operation and object tracking? (Choose two.)

- A. Configure an IP SLA operation to test the primary path.
- B. Associate an object-tracking statement with the IP SLA operation.
- C. Configure equal administrative distances on the primary and backup static routes.
- D. Apply the tracking object to the backup route only.
- E. Configure BFD under the static route to create the tracking object automatically.

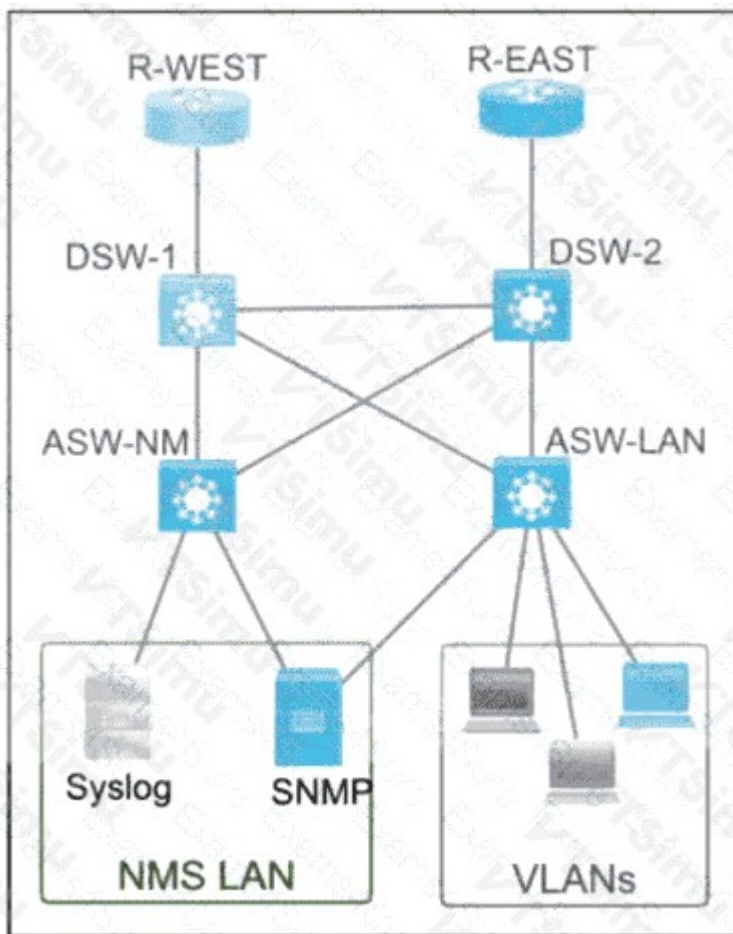
ANSWER: A B

Explanation:

An IP SLA operation must be configured to test the reachability of the primary path. For example, an ICMP echo operation can probe a remote address through the preferred next hop at a defined frequency. An object-tracking statement then tracks the IP SLA operation, allowing route installation to depend on the probe state.

The primary static route is configured with the tracking object, and the backup static route is configured with a higher administrative distance. When the IP SLA operation fails, the tracked primary route is removed from the routing table and the floating static route becomes active. When the operation returns to an up state, the primary route is reinstalled. Cisco documents IP SLA tracking and static-route integration in the [Cisco IOS XE IP SLAs and Object Tracking Configuration Guide](#).

QUESTION NO: 46 - (SIMULATION)



Troubleshoot R-WEST to achieve the desired results:

1. The locally generated logs should have sequence numbers, date and time.
2. The SNMP traps related to OSPF and participating interface state changes utilizing RFC1253-MIB OSPFv2 should be send to SNMP server.

ANSWER: See the explanation for the answer

Explanation:

On R-WEST, enable sequence numbers and date/time stamps for local log messages, then enable the standard OSPF SNMP traps. The plain `ospf` trap option enables the RFC 1253 OSPFv2 MIB notifications, including OSPF interface state-change notifications.

```
enable
configure terminal
service sequence-numbers
service timestamps log datetime msec
snmp-server enable traps ospf
end
write memory
```

Do **not** use only the Cisco-specific OSPF trap option for this requirement: the requested OSPFv2 interface-state notification is from the standard RFC 1253 MIB. Assuming the SNMP server destination/community is already configured in the supplied lab, these commands cause the enabled traps to be sent to that SNMP server.

QUESTION NO: 47

Which two statements describe IPsec protection for a DMVPN tunnel? (Choose two.)

- A. An IPsec profile can be applied to the tunnel interface.
- B. GRE encapsulation provides encryption for NHRP and user traffic.
- C. IPsec provides confidentiality and integrity for protected tunnel traffic.
- D. IPsec eliminates the need for NHRP mappings in a multipoint GRE deployment.
- E. IPsec requires one point-to-point GRE tunnel interface for every spoke.

ANSWER: A C

Explanation:

DMVPN commonly applies IPsec protection directly to the GRE tunnel interface by using the `tunnel protection ipsec profile` command. The profile associates the tunnel with an IPsec transform set or IPsec proposal, depending on the software release and configuration model. Therefore, **an IPsec profile can be applied to the tunnel interface** is correct.

GRE provides multipoint tunnel encapsulation and supports the transport of multicast and routing-protocol packets, but GRE does not encrypt traffic. IPsec supplies the confidentiality, integrity, and peer authentication functions required to protect the DMVPN overlay. Cisco documents tunnel protection for DMVPN in the [Cisco IOS XE DMVPN Configuration Guide](#).

QUESTION NO: 48

What are two functions of IPv6 Source Guard? (Choose two.)

- A. It uses the populated binding table for allowing legitimate traffic.
- B. It works independent from IPv6 neighbor discovery.
- C. It denies traffic from unknown sources or unallocated addresses.
- D. It denies traffic by inspecting neighbor discovery packets for specific pattern.
- E. It blocks certain traffic by inspecting DHCP packets for specific sources.

ANSWER: A C

Explanation:

IPv6 Source Guard validates the source IPv6 address of traffic received on an enabled Layer 2 access port. It uses the IPv6 binding table to determine which source addresses are authorized on that port, so "It uses the populated binding table for allowing legitimate traffic" is correct. The binding information can be learned through IPv6 Snooping mechanisms, including DHCPv6 snooping and Neighbor Discovery inspection, or can be configured statically. This port-specific validation prevents a host from transmitting using an address that is not associated with its valid binding.

Accordingly, "It denies traffic from unknown sources or unallocated addresses" is also correct. When a packet's source IPv6 address has no matching authorized binding for the ingress interface, IPv6 Source Guard drops that traffic. This enforcement helps mitigate IPv6 address spoofing on access networks and ensures that hosts use only addresses assigned or validated for their switch port. IPv6 Source Guard is part of Cisco IPv6 First-Hop Security, which combines binding-based enforcement with other controls to protect hosts and the local IPv6 segment. Cisco documents the feature in its [IPv6 First-Hop Security Configuration Guide](#); see also Cisco's [IPv6 First-Hop Security overview](#).

Exhibit.

```
R1# show route-map
route-map Redistribution_EIGRP, permit, sequence 10
Match clauses:
  ip address (access-lists): 10
Set clauses:
  tag 666
Policy routing matches: 0 packets, 0 bytes
route-map Redistribution_EIGRP, permit, sequence 20
Match clauses:
Set clauses:
Policy routing matches: 0 packets, 0 bytes

R1# show access-lists
Standard IP access list 10
  10 permit 172.16.1.0, wildcard bits 0.0.0.255
  20 permit 172.16.2.0, wildcard bits 0.0.0.255
```

Refer to the exhibit. The router is redistributing a prefix 172.16.10.0/24 that should have been filtered. Which action resolves the issue?

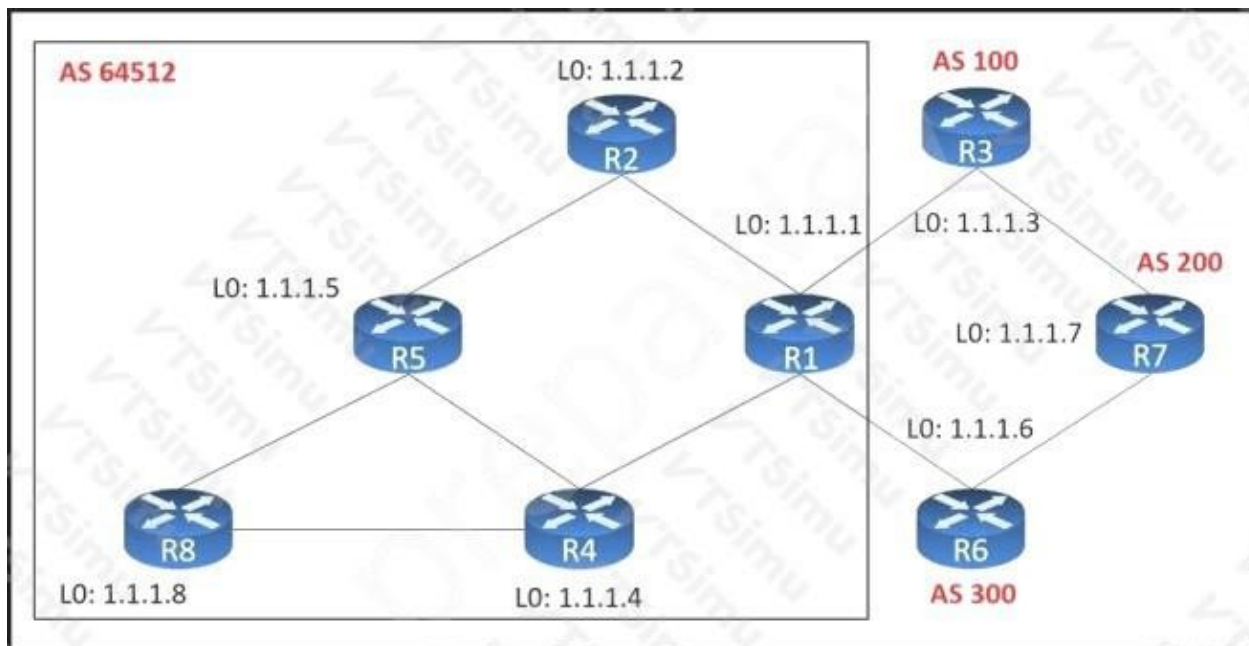
- A. Permit the route in route-map sequence 20.
- B. Add the route in access-list 10.
- C. Match the tag 666 for the route in the route map.
- D. Remove route-map sequence 20.

ANSWER: D

Explanation:

Remove route-map sequence 20. In a route map used for redistribution, entries are evaluated in sequence order. The filtering entry is intended to prevent the identified prefix from being redistributed. However, the later permit entry has no match conditions, so it acts as a catch-all: any route that does not match an earlier route-map entry is permitted by that sequence. This includes 172.16.10.0/24 when the access-list logic causes it not to satisfy the preceding route-map match. Removing the unconditional permit sequence restores the route map's implicit deny at the end. Consequently, routes that are not explicitly permitted by an earlier route-map sequence, including the unwanted prefix, are not redistributed. This is a common route-filtering design consideration: an empty `permit` route-map clause permits every route that reaches it, rather than only routes intended for redistribution. Cisco documents that route-map clauses are processed in order and that route processing ends according to the matching permit or deny result. See Cisco's [route redistribution with route maps guidance](#) and the [Cisco IOS route-policy documentation](#).

QUESTION NO: 50



Refer to the exhibit. An engineer configured R2 and R5 as route reflectors and noticed that not all routes are sent to R1 to advertise to the eBGP peers.

Which iBGP routers must be configured as route reflectors to advertise all routes to restore reachability across all networks?

- A. R1 and R4
- B. R1 and R5
- C. R4 and R5
- D. R2 and R5

ANSWER: C

Explanation:

R4 and R5 must act as the route reflectors. In an iBGP design, a normal BGP speaker does not advertise routes learned from one iBGP peer to another iBGP peer. Route reflection removes the need for a full iBGP mesh by permitting a route reflector to advertise eligible routes between its clients and to other BGP peers while preserving loop-prevention information such as the ORIGINATOR_ID and CLUSTER_LIST attributes.

Based on the topology shown, R4 and R5 are the devices positioned to reflect the routes from the relevant iBGP portions of the network so that R1 receives the complete set of internal BGP prefixes. Once R1 learns those prefixes, it can advertise the appropriate routes to its eBGP neighbors, restoring end-to-end reachability. The route-reflector client relationships must also be configured consistently on the affected peers, typically using the `neighbor ... route-reflector-client` command under the BGP address family. This design retains iBGP loop prevention while ensuring that route propagation reaches the eBGP edge router.

Cisco's overview of BGP route reflection is available in [Cisco BGP Route Reflector documentation](#). The route-reflection rules and loop-prevention attributes are defined in [RFC 4456](#).

QUESTION NO: 51

Refer to the exhibit.

```
*Sep 26 19:50:43.504: SNMP: Packet received via UDP from
192.168.1.2 on GigabitEthernet0/1SrParseV3SnmpMessage: No
matching Engine ID.
```

```
SrParseV3SnmpMessage: Failed.
```

```
SrDoSnmp: authentication failure, Unknown Engine ID
```

```
*Sep 26 19:50:43.504: SNMP: Report, reqid 29548, errstat 0,
erridx 0
```

```
internet.6.3.15.1.1.4.0 = 3
```

```
*Sep 26 19:50:43.508: SNMP: Packet sent via UDP to 192.168.1.2
process_mgmt_req_int: UDP packet being de-queued
```

Which two commands provide the administrator with the information needed to resolve the issue? (Choose two.)

- A. debug snmpv3 engine-id
- B. show snmp user
- C. debug snmp packet
- D. debug snmp engine-id
- E. show snmpv3 user

ANSWER: B C

Explanation:

`show snmp user` provides the locally configured SNMPv3 user information needed to validate the device-side SNMPv3 security configuration. Its output identifies configured users and their security model, authentication protocol, privacy protocol, group association, and engine-ID-related details. This allows the administrator to confirm that the user expected by the SNMP manager is present and that the device has the necessary SNMPv3 user credentials and security settings available for the requested access.

`debug snmp packet` supplies real-time visibility into SNMP packet processing on the device. During an attempted poll, it can reveal the SNMP request and response activity and expose details relevant to an SNMPv3 exchange, such as the security level used and errors encountered while processing the message. Used together, the configured-user view and packet-level debug distinguish configuration state from the behavior of the actual transaction, providing the information required to troubleshoot an SNMPv3 communication problem. Cisco documents SNMP user verification and SNMP debugging commands in the [Cisco IOS XE SNMP Command Reference](#) and SNMPv3 configuration guidance in the [Cisco IOS XE SNMPv3 Configuration Guide](#).

QUESTION NO: 52

Users were moved from the local DHCP server to the remote corporate DHCP server. After the move, none of the users were able to use the network.

Which two issues will prevent this setup from working properly? (Choose two)

- A. Auto-QoS is blocking DHCP traffic.
- B. The DHCP server IP address configuration is missing locally

- C. 802.1X is blocking DHCP traffic
- D. The broadcast domain is too large for proper DHCP propagation
- E. The route to the new DHCP server is missing

ANSWER: B E

Explanation:

A remote DHCP server cannot receive a client's initial DHCPDISCOVER broadcast unless the Layer 3 gateway for the client VLAN is configured as a DHCP relay. The local gateway interface must contain an `ip helper-address` statement identifying the corporate DHCP server. Therefore, "The DHCP server IP address configuration is missing locally" represents the missing local relay-agent configuration required to convert the client broadcast into a unicast DHCP relay message. The relay also places gateway-address information in the request so that the central DHCP server can select the correct scope for the originating subnet.

"The route to the new DHCP server is missing" also prevents operation because DHCP relay depends on bidirectional IP connectivity. The relay device must be able to route its unicast request to the DHCP server, and the server must have a return route to the relay/client subnet so that the DHCPOFFER and subsequent DHCP messages can be delivered back through the relay. After moving DHCP services centrally, validating the helper address, routing table, and return-path routing is essential for every affected VLAN. Cisco documents the DHCP relay process and the use of `ip helper-address` in its [DHCP configuration guidance](#); the protocol behavior is also defined in [RFC 2131](#).

QUESTION NO: 53

A newly Installed router starts establishing an LDP session from another MPLS router to which it is not directly connected. Which LDP message type responds by target router to the Initiating router using UDP protocol?

- A. notification message
- B. session message
- C. extended discovery (targeted Hello) message
- D. advertisement message

ANSWER: C

Explanation:

The **extended discovery (targeted Hello) message** is correct because LDP uses extended discovery to discover and establish an LDP relationship between LSRs that are not directly connected. This process uses targeted Hello messages, which are sent as UDP packets to port 646. The initiating router sends a targeted Hello to the remote router's transport address, and the target router can respond with its own targeted Hello. These exchanged Hellos communicate the peers' LDP identifiers, transport addresses, hold times, and relevant discovery information needed before an LDP session can be formed.

After the targeted discovery exchange establishes that both routers are willing to peer, the LDP session itself is established over TCP, also using port 646. Therefore, UDP is specifically associated with the discovery phase, while TCP provides the reliable transport for LDP session establishment, label advertisement, notification, and keepalive-related session communication. In this scenario, the routers are nonadjacent, so the discovery mechanism is extended discovery rather than link-local basic discovery. RFC 5036 defines LDP discovery and specifies the use of targeted Hellos for extended discovery: [RFC 5036: LDP Specification](#). Cisco also describes targeted LDP as the mechanism for establishing LDP adjacencies with non-directly connected peers: [Cisco MPLS LDP Discovery Configuration Guide](#).

QUESTION NO: 54

Refer to the exhibit.

```
ip dhcp pool 1
network 200.30.30.0/24
default-router 200.30.30.100
lease 40
!
ip dhcp pool 2
network 200.30.40.0/24
default-router 200.30.40.100
lease 40
!
```

The server for the finance department is not reachable consistently on the 200.30.40.0/24 network and after every second month it gets a new IP address. Which two actions must be taken to resolve this issue? (Choose two.)

- A. Configure the server to use DHCP on the network with default gateway 200.30.40.100.
- B. Configure the server with a static IP address and default gateway.
- C. Configure the router to exclude a server IP address.
- D. Configure the server to use DHCP on the network with default gateway 200.30.30.100.
- E. Configure the router to exclude a server IP address and default gateway.

ANSWER: B C

Explanation:

Configuring the server with a static IP address and default gateway gives the finance server a permanent, predictable address in the 200.30.40.0/24 subnet. A server that provides a departmental service should retain a stable address so that clients, DNS records, access policies, and any configured routes continue to point to the same destination. Its default gateway must be the router interface in that subnet so replies to clients on remote networks are routed correctly.

Configuring the router to exclude a server IP address reserves the selected static address from the router's DHCP allocation process. This prevents the DHCP server from leasing the server's manually assigned address to another endpoint, which would create an IP-address conflict and intermittent reachability. The excluded address must be chosen from the DHCP pool's subnet but omitted from dynamic assignment. Together, static addressing on the server and DHCP exclusion on the router ensure that the address remains both stable and unique.

Cisco documents DHCP exclusions with the `ip dhcp excluded-address` command and explains DHCP pool configuration in its [IOS XE DHCP Server Configuration Guide](#). Cisco's guidance on static addressing and DHCP reservations is also covered in the [Cisco DHCP configuration documentation](#).

QUESTION NO: 55

Which two statements about VRF-Lite configurations are true? (Choose two.)

- A. They support the exchange of MPLS labels
- B. Different customers can have overlapping IP addresses on different VPNs
- C. They support a maximum of 512,000 routes
- D. Each customer has its own dedicated TCAM resources
- E. Each customer has its own private routing table.
- F. They support IS-IS

Explanation:

VRF-Lite provides virtual routing and forwarding instances on a router or multilayer switch without requiring an MPLS provider core. A separate VRF is associated with the applicable customer-facing interfaces, causing the device to maintain an independent routing and forwarding context for that customer or tenant. Therefore, "Each customer has its own private routing table." is correct: routes learned through static configuration or supported routing protocols are installed in the routing table for the relevant VRF rather than the global routing table or another customer's VRF. VRF-specific forwarding information is then derived from that isolated routing context.

"Different customers can have overlapping IP addresses on different VPNs" is also correct. Because route lookup is performed within the selected VRF, the same IPv4 or IPv6 prefix can exist in multiple VRFs without ambiguity. This separation is a principal VRF-Lite use case, allowing a shared enterprise infrastructure to support isolated customer, department, or tenant address spaces. Controlled communication between VRFs, if required, is configured explicitly through route leaking or another designed inter-VRF service; isolation is the normal behavior. Cisco's VRF-Lite configuration documentation describes the use of separate routing tables and interface-to-VRF assignment: [Cisco IOS VRF Configuration Guide](#). Cisco also summarizes VRF technology and its routing-table separation model at [Cisco VRF overview](#).

QUESTION NO: 56

Refer to the exhibit.

```
Router#show access-lists
Standard IP access list 1
    10 permit 192.168.2.2 (1 match)
Router#
Router#show route-map
route-map RM-OSPF-DL, deny, sequence 10
  Match clauses:
    ip address (access-lists): 1
  Set clauses:
  Policy routing matches: 0 packets, 0 bytes
Router#
Router#show running-config | section ospf
router ospf 1
  network 192.168.1.1 0.0.0.0 area 0
  network 192.168.12.0 0.0.0.255 area 0
  distribute-list route-map RM-OSPF-DL in
Router#
```

Which two actions should be taken to access the server? (Choose two.)

- A. Modify the access list to add a second line of permit ip any
- B. Modify the access list to deny the route to 192.168.2.2.

- C. Modify distribute list seq 10 to permit the route to 192.168.2.2.
- D. Add a sequence 20 in the route map to permit access list 1.
- E. Add a floating static route to reach to 192.168.2.2 with administrative distance higher than OSPF

ANSWER: B E

Explanation:

“Modify the access list to deny the route to 192.168.2.2.” is required to prevent that specific server route from being accepted by the routing-policy match used in the exhibited redistribution or route-filtering process. In a route-map policy, the access list controls whether a prefix matches the route-map sequence. Denying the host route removes it from the affected routing exchange, ensuring that the undesired OSPF-learned path is not selected for this destination. Cisco documents that route maps can be used to control which routes are redistributed and that match conditions determine the routes to which the policy applies; see [Cisco’s OSPF redistribution documentation](#).

“Add a floating static route to reach to 192.168.2.2 with administrative distance higher than OSPF” provides the required alternate path. A floating static route has an administrative distance deliberately set higher than the primary dynamic routing protocol, so it remains installed as a standby route and becomes active only when the preferred route is unavailable. This preserves normal OSPF route selection while ensuring continued reachability to the server after the route-filtering change or a primary-path failure. Cisco’s [administrative-distance overview](#) describes how Cisco routers select routes based on administrative distance.

QUESTION NO: 57

Refer to the exhibit.

```
snmp-server community ciscotest1
snmp-server host 192.168.1.128 ciscotest
snmp-server enable traps bgp
```

Network operations cannot read or write any configuration on the device with this configuration from the operations subnet. Which two configurations fix the issue? (Choose two.)

- A. Configure SNMP rw permission in addition to community ciscotest.
- B. Modify access list 1 and allow operations subnet in the access list.
- C. Modify access list 1 and allow SNMP in the access list.
- D. Configure SNMP rw permission in addition to version 1.
- E. Configure SNMP rw permission for community ciscotest with access list 1.

ANSWER: B E

Explanation:

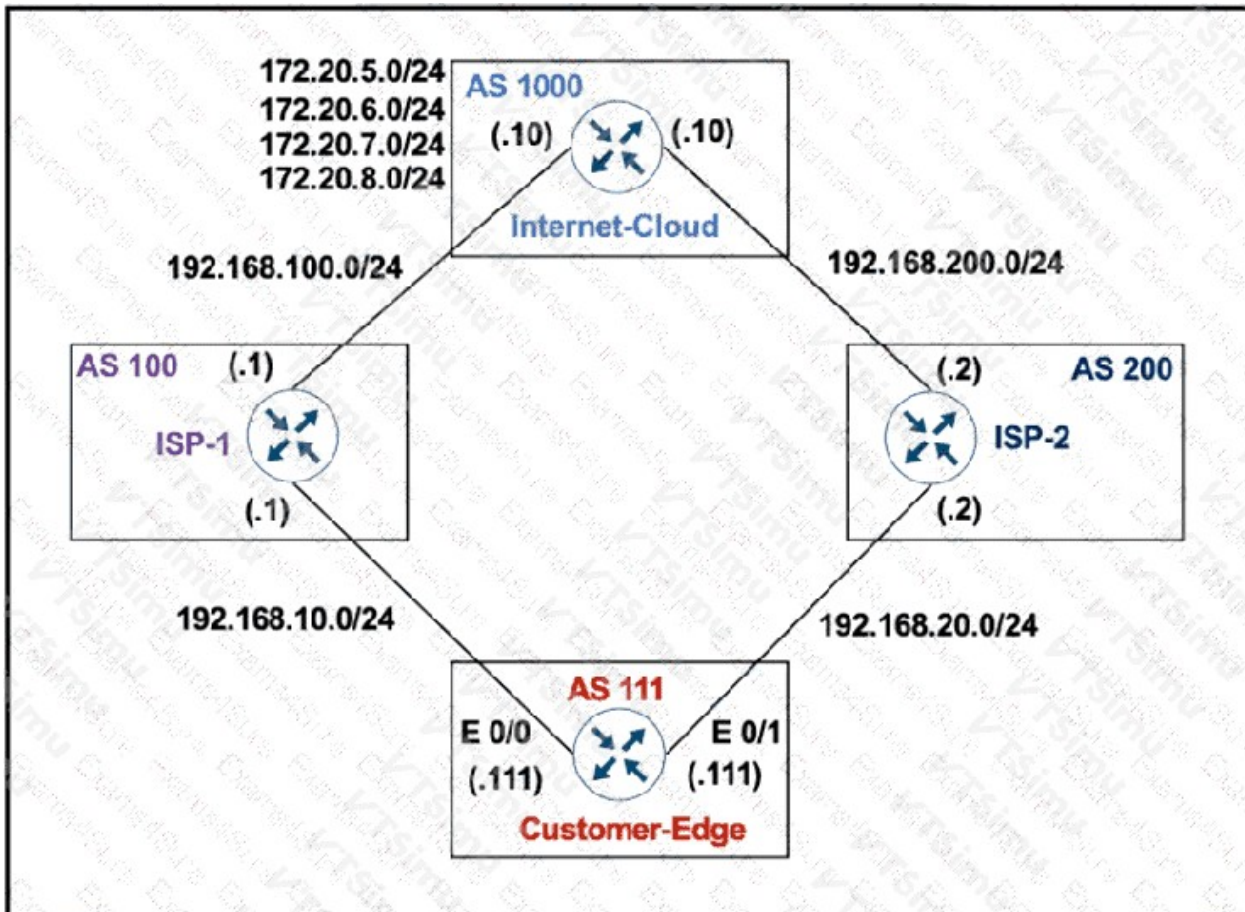
The required corrections are to permit the operations subnet in access list 1 and to configure the `ciscotest` SNMP community with read-write access while retaining access-list association 1. In an SNMP community configuration, the ACL specified after the community string restricts which source addresses may use that community. Therefore, access list 1 must explicitly permit the network-management hosts in the operations subnet; otherwise, requests from those hosts are rejected before SNMP read or write access can be granted. The ACL is evaluated against the source IP address of the SNMP manager, so its permit statement must match the operations subnet.

Configuration changes through SNMP also require read-write authorization. A community configured as `ro` can retrieve permitted management information but cannot perform SNMP SET operations, which are needed to modify writable configuration objects. Configuring `snmp-server community ciscotest rw 1` provides write capability and continues to limit use of that community to the sources permitted by ACL 1. This pairing follows Cisco’s community-based SNMP

access-control model: community privilege controls read-only versus read-write operations, and the referenced ACL controls eligible manager addresses. See Cisco's [IOS XE SNMP Configuration Guide](#) and the [Cisco IOS SNMP Command Reference](#).

QUESTION NO: 58

Refer to Exhibit:



Customer-Edge

```
ip prefix-list PLIST1 permit 172.20.5.0/24
!
route-map SETLP permit 10
  match ip address prefix-list PLIST1
  set local-preference 90
!
router bgp 111
  neighbor 192.168.10.1 remote-as 100
  neighbor 192.168.10.1 route-map SETLP in
  neighbor 192.168.20.2 remote-as 200
```

AS 111 wanted to use AS 200 as the preferred path for 172.20.5.0/24 and AS 100 as the backup. After the configuration, AS 100 is not used for any other routes. Which configuration resolves the issue?

- A. route-map SETLP permit 10 match ip address prefix-list PLIST1 set local-preference 99 route-map SETLP permit 20
- B. route-map SETLP permit 10 match ip address prefix-list PLIST1 set local-preference 110 route-map SETLP permit 20
- C. router bgp 111 no neighbor 192.168.10.1 route-map SETLP in neighbor 192.168.10.1 route-map SETLP out
- D. router bgp 111 no neighbor 192.168.10.1 route-map SETLP in neighbor 192.168.20.2 route-map SETLP in

ANSWER: A

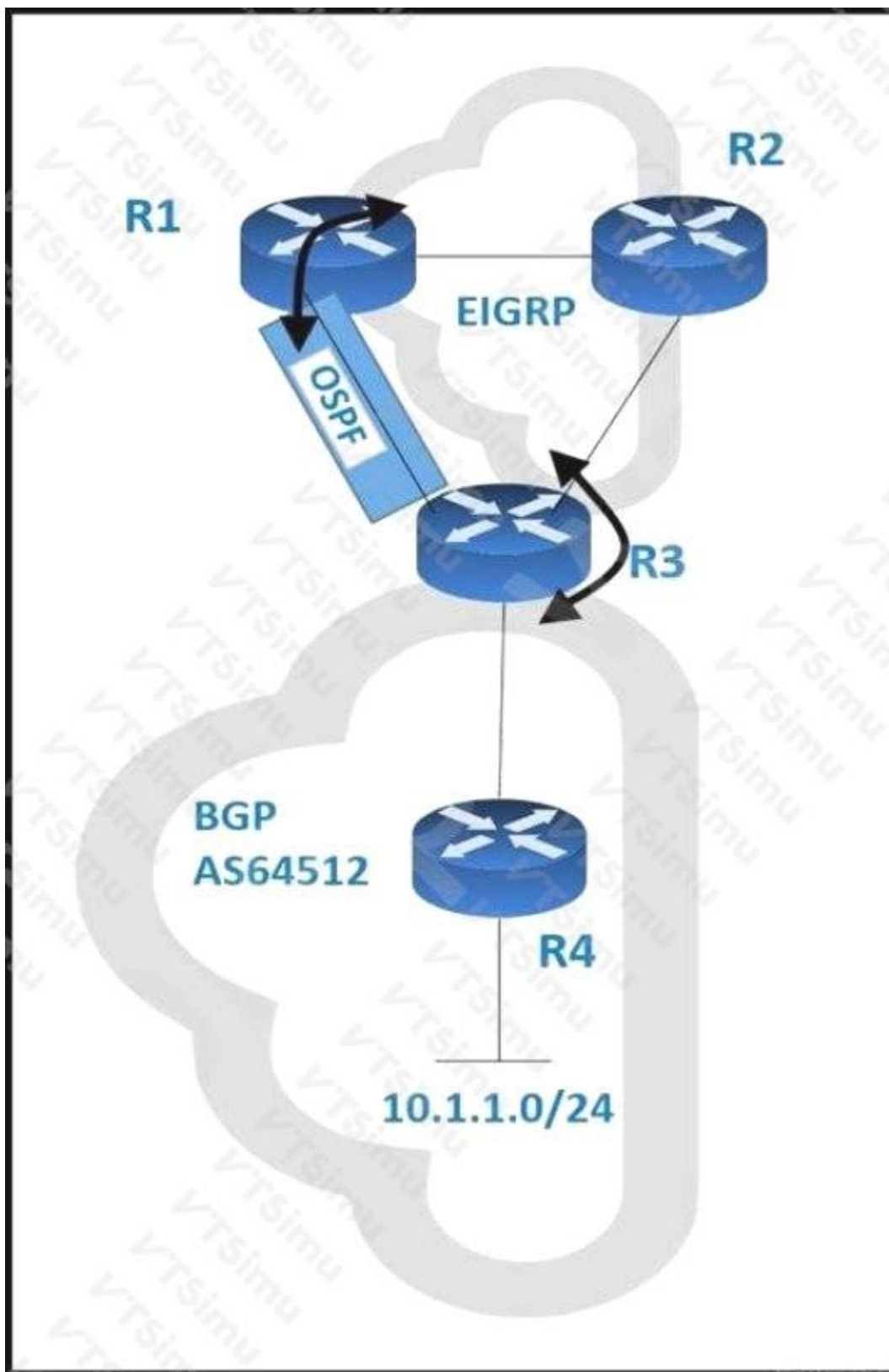
Explanation:

The configuration `route-map SETLP permit 10 match ip address prefix-list PLIST1 set local-preference 99 route-map SETLP permit 20` resolves the issue because it handles both the desired backup policy and the remaining routes received from AS 100. Local preference is a BGP attribute used within an autonomous system, and BGP selects the path with the higher local-preference value. Assigning a local preference of 99 only to the prefix matched by `PLIST1` makes the route through AS 100 less preferred than the corresponding route through AS 200 when AS 200 retains the normal local preference of 100. This produces the required primary/backup behavior for 172.20.5.0/24.

The additional `permit 20` sequence is essential because route maps have an implicit deny when no sequence matches. It permits all other routes received from AS 100 without modifying their local preference, so they remain eligible BGP paths at the normal default local-preference value. Cisco documents local preference as a well-known discretionary BGP attribute used to influence outbound path selection from an AS, and its BGP best-path process prefers the highest value. See [Cisco BGP Best Path Selection Algorithm](#) and [Cisco BGP Path Selection documentation](#).

QUESTION NO: 59

Refer to the exhibit.



BGP and EIGRP are mutually redistributed on R3, and EIGRP and OSPF are mutually redistributed on R1. Users report packet loss and interruption of service to applications hosted on the 10.1.1.0/24 prefix. An engineer tested the link from R3 to R4 with no packet loss present but has noticed frequent routing changes on R3 when running the debug ip route command.

Which action stabilizes the service?

- A. Reduce frequent OSPF SPF calculations on R3 that cause a high CPU and packet loss on traffic traversing R3.
- B. Tag the 10.1.1.0/24 prefix and deny the prefix from being redistributed into OSPF on R1.
- C. Place an OSPF distribute-list outbound on R3 to block the 10.1.1.0/24 prefix from being advertised back to R3.

D. Repeat the test from R4 using ICMP ping on the local 10.1.1.0/24 prefix, and fix any Layer 2 errors on the host or switch side of the subnet.

ANSWER: B

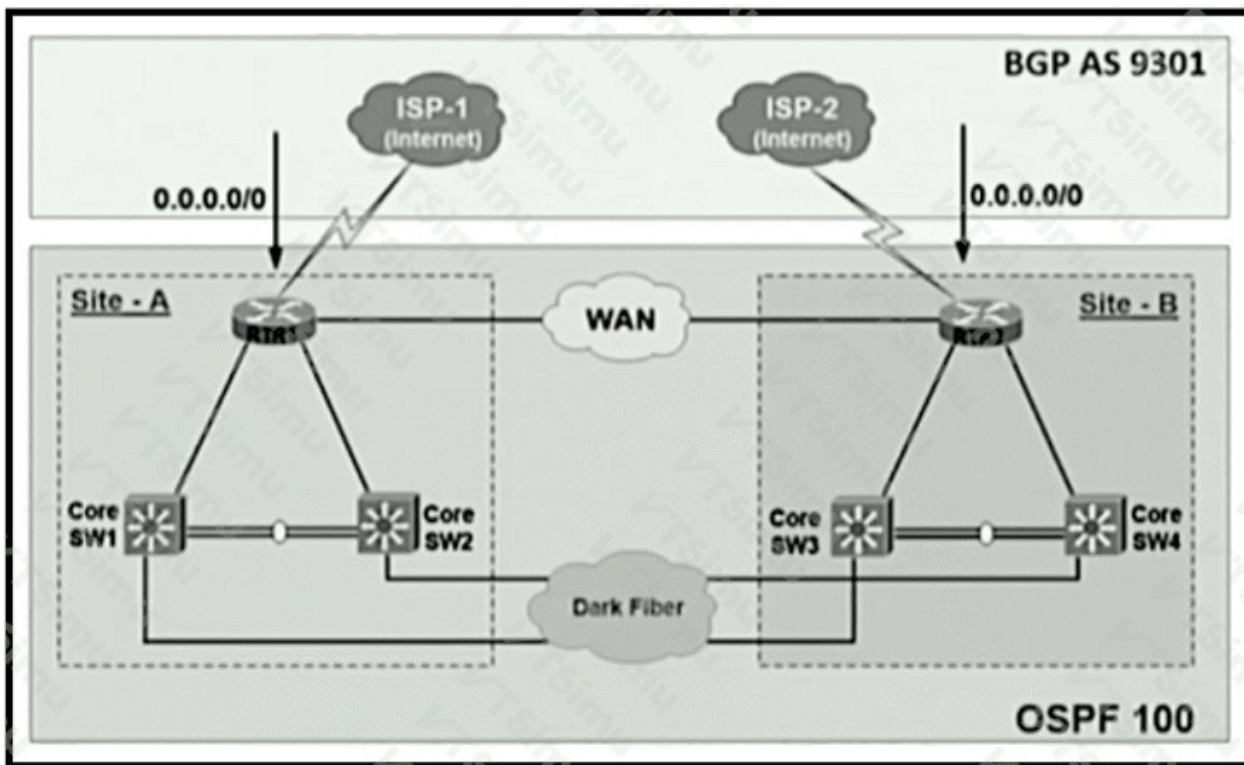
Explanation:

Tag the 10.1.1.0/24 prefix and deny the prefix from being redistributed into OSPF on R1. This prevents a route-redistribution feedback loop. The application subnet is learned through BGP at R3 and is redistributed into EIGRP. Because EIGRP is then redistributed into OSPF at R1, that same reachability information can circulate through the multiple routing domains and return as a competing route. Repeated installation and withdrawal of those redistributed routes causes the frequent route changes observed on R3, resulting in intermittent forwarding changes and application disruption.

Route tags provide a persistent attribute for identifying routes that originated from a particular redistribution boundary. The route should be tagged when it enters the redistribution path, and R1 should use a route map that matches that tag and prevents the tagged route from being injected into OSPF. This preserves valid routing toward 10.1.1.0/24 while ensuring that the route cannot re-enter another protocol domain and feed back toward its origin. Cisco recommends route tagging and route maps to control redistribution and avoid routing loops. See Cisco's [Redistributing Routing Protocols](#) guidance and the [OSPF Version 2 specification](#) for external-route tag behavior.

QUESTION NO: 60

Refer to the exhibit.



The Internet traffic should always prefer Site-A ISP-1 if the link and BGP connection are up; otherwise, all Internet traffic should go to ISP-2. Redistribution is configured between BGP and OSPF routing protocols and it is not working as expected. What action resolves the issue?

- A. Set metric-type 2 at Site-A RTR1, and set metric-type 1 at Site-B RTR2
- B. Set OSPF cost 100 at Site-A RTR1, and set OSPF Cost 200 at Site-B RTR2
- C. Set OSPF cost 200 at Site: A RTR1 and set OSPF Cost 100 at Site-B RTR2
- D. Set metric-type 1 at Site-A RTR1, and set metric-type 2 at Site-B RTR2

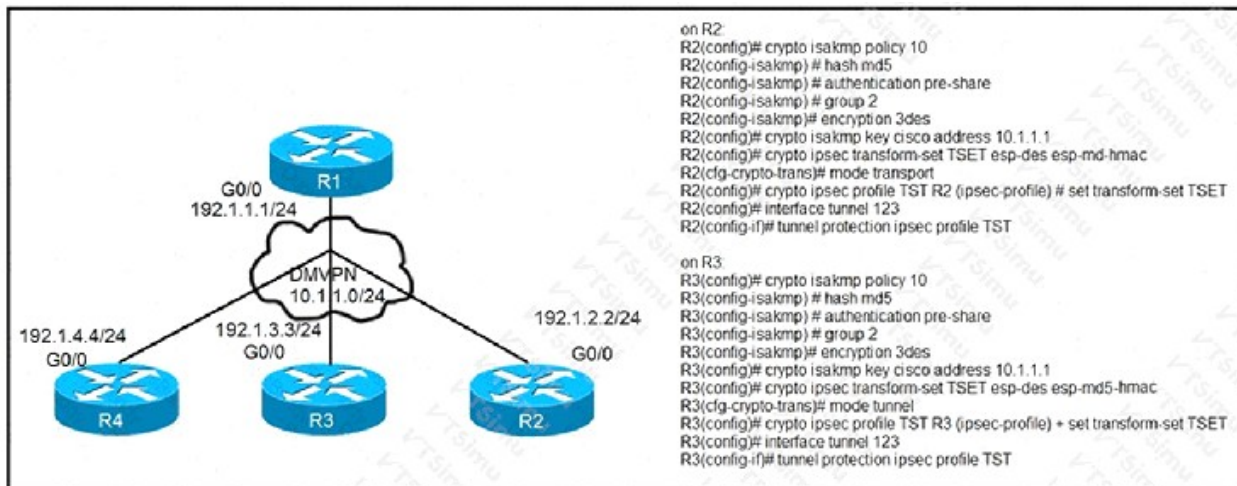
ANSWER: D

Explanation:

Set metric-type 1 at Site-A RTR1, and set metric-type 2 at Site-B RTR2. BGP-learned Internet routes redistributed into OSPF become OSPF external routes. OSPF selects external type 1 routes before external type 2 routes, provided routes being compared have otherwise reached the external-route selection stage. Therefore, advertising the Site-A ISP-1 routes as type 1 external routes makes that exit point preferred throughout the OSPF domain whenever the Site-A BGP-learned routes are available. This preference does not depend on which internal router is closer to either autonomous system boundary router: an external type 1 route is preferred over an external type 2 route during OSPF route selection. If the Site-A link or BGP adjacency fails and its redistributed routes are withdrawn, the type 1 routes disappear from OSPF. The remaining type 2 routes redistributed by Site-B RTR2 are then selected, providing the intended ISP-2 backup path. Cisco documents that type 1 external metrics include the internal cost to the advertising router, while type 2 routes use their external metric for comparison; OSPF also gives type 1 external paths preference over type 2 paths. See [Cisco: OSPF Design Guide](#) and [RFC 2328, OSPF external-route calculation](#).

QUESTION NO: 61

Refer to the exhibit.



After applying IPsec, the engineer observed that the DMVPN tunnel went down, and both spoke-to-spoke and hub were not establishing. Which two actions resolve the issue? (Choose two.)

- A. Configure the crypto isakmp key cisco address 192.1.1.1 on R2 and R3
- B. Configure the crypto isakmp key cisco address 0.0.0.0 on R2 and R3.
- C. Change the mode from mode tunnel to mode transport on R3
- D. Change the mode from mode transport to mode tunnel on R2.
- E. Remove the crypto isakmp key cisco address 10.1.1.1 on R2 and R3

ANSWER: B C

Explanation:

Configuring the ISAKMP preshared key with the peer address 0.0.0.0 on both spokes allows either spoke to authenticate IPsec negotiations from dynamically discovered DMVPN peers. This is essential to support dynamic spoke-to-spoke tunnels, because the remote NBMA peer address is not known in advance in the same way as a fixed site-to-site IPsec peer. The wildcard preshared key also permits the spokes to establish the required security associations when NHRP redirects traffic directly between them.

Changing the IPsec transform-set on R3 from tunnel mode to transport mode makes the IPsec encapsulation consistent with the DMVPN design. DMVPN runs GRE over an NBMA network, and GRE already supplies the tunnel encapsulation needed

for the overlay. Applying IPsec transport mode protects the GRE packet while retaining the GRE header information required by DMVPN and NHRP. Both ends participating in the protected DMVPN overlay must use compatible IPsec settings, including the transform-set mode, for IPsec security associations and the DMVPN tunnel to establish successfully. Cisco documents DMVPN IPsec protection and the use of transport mode for GRE/DMVPN in its [DMVPN IPsec configuration guide](#) and [IKE troubleshooting documentation](#).

QUESTION NO: 62

Exhibit:

```
policy-map COPP-7600
class COPP-CRITICAL-7600
  police cir 2000000 bc 62500
  conform-action transmit
  exceed-action transmit
!
class class-default
  police cir 200000 bc 6250
  conform-action transmit
  exceed-action drop
!
class-map match-all COPP-CRITICAL-7600
  match access-group name COPP-CRITICAL-7600
!
ip access-list extended COPP-CRITICAL-7600
  permit ip any any eq http
  permit ip any any eq https
```

BGP is flapping after the Copp policy is applied. What are the two solutions to fix the issue?

(Choose two)

- A. Configure BGP in the COPP-CRITICAL-7600 ACL
- B. Configure a higher value for CIR under the default class to allow more packets during peak traffic
- C. Configure a higher value for CIR under the class COPP-CRITICAL-7600
- D. Configure a three-color policer instead of two-color policer under class COPP-CRITICAL-7600
- E. Configure IP CEF to CoPP policy and BGP to work

ANSWER: A B

Explanation:

Configuring BGP in the COPP-CRITICAL-7600 ACL is a valid solution because CoPP classifies control-plane traffic using the ACL associated with the class map. Adding a permit entry that matches BGP TCP traffic, typically TCP destination port 179, places BGP packets into the explicitly defined critical CoPP class rather than allowing them to fall through to `class-default`. This gives BGP traffic the treatment intended for protected, important control-plane protocols. BGP uses TCP port 179 for its transport connection, as specified in [RFC 4271](#).

Configuring a higher value for CIR under the default class to allow more packets during peak traffic is also valid. In the displayed policy, traffic not matched by the critical ACL, including the currently unclassified BGP packets, is handled by the

default class policer. If legitimate control-plane traffic exceeds that policer's committed rate, excess packets are dropped. Raising the default-class CIR reduces those drops and can stabilize the BGP session during traffic bursts. CoPP is specifically intended to classify and police traffic destined to the control plane, so the policer rates must accommodate expected legitimate protocol traffic; Cisco describes this behavior in its [Control Plane Policing documentation](#).

QUESTION NO: 63

Refer to the exhibit.

The R2 loopback interface is advertised with RIP and EIGRP using default values. Which configuration changes make R1 reach the R2 loopback using RIP?

- A. R1(config)# router ripR1(config-router)# distance 90
- B. R1(config)# router ripR1(config-router)# distance 100
- C. R1(config)# router eigrp 1R1(config-router)# distance eigrp 130 120
- D. R1(config)# router eigrp 1R1(config-router)# distance eigrp 120 120
- E. R1(config)# router eigrp 1
R1(config-router)# distance eigrp 130 121

ANSWER: E

Explanation:

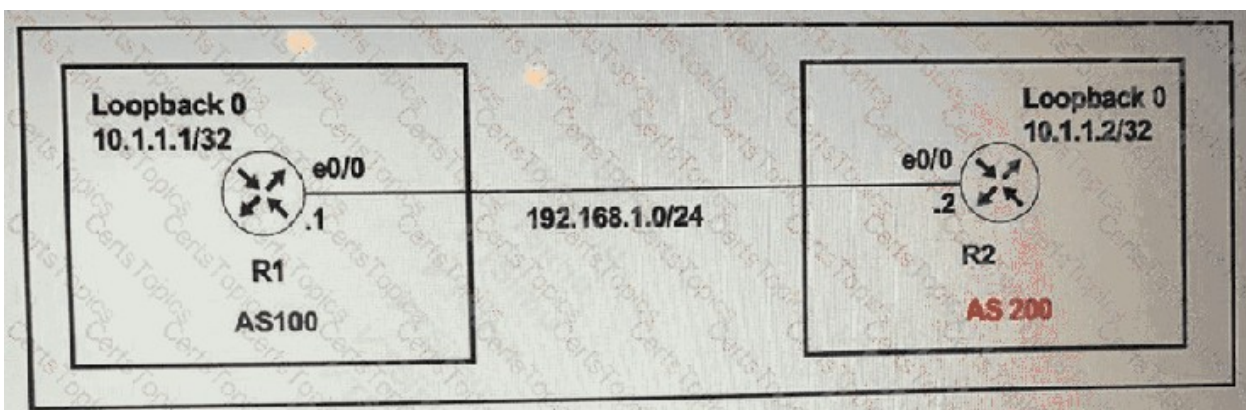
```
R1(config)# router eigrp 1
```

R1(config-router)# distance eigrp 130 121 is correct because it raises the administrative distance for internally learned EIGRP routes to 121. By default, RIP routes have an administrative distance of 120, while internal EIGRP routes have an administrative distance of 90. When R1 learns the same R2 loopback prefix through both protocols, it selects the route with the lower administrative distance before evaluating protocol-specific metrics. Setting the EIGRP internal distance to 121 makes the RIP route's default distance of 120 lower, so R1 installs and uses the RIP-learned route to reach the loopback interface.

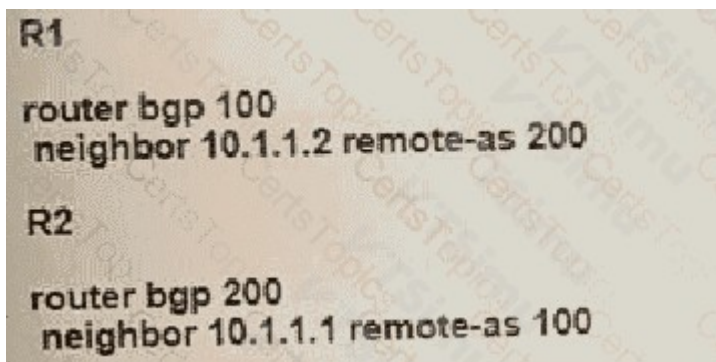
The `distance eigrp` command takes the external EIGRP administrative distance first and the internal EIGRP administrative distance second. Retaining 130 for external EIGRP routes is appropriate here; changing the internal value to 121 is the required change because the advertised loopback is learned as an internal EIGRP route. Cisco documents the default administrative-distance values at [Cisco EIGRP documentation](#) and the EIGRP distance command syntax at [Cisco IOS EIGRP command reference](#).

QUESTION NO: 64

Refer to the exhibit.



The R1 and R2 configurations are:



The neighbor is not coming up. Which two sets of configurations bring the neighbors up? (Choose two.)

- A. R1ip route 10.1.1.1 255.255.255.255 192.168.1.1! router bgp 200neighbor 10.1.1.1 ttl-security hops 1neighbor 10.1.1.1 update-source loopback 0
- B. R2ip route 10.1.1.1 255.255.255.255 192.168.1.1! router bgp 200neighbor 10.1.1.1 disable-connected-checkneighbor 10.1.1.1 update-source loopback 0
- C. R2ip route 10.1.1.2 255.255.255.255 192.168.1.2! router bgp 100neighbor 10.1.1.2 ttl-security hops 1neighbor 10.1.1.2 update-source loopback 0
- D. R1ip route 10.1.1.2 255.255.255.255 192.168.1.2! router bgp 100neighbor 10.1.1.1 ttl-security hops 1neighbor 10.1.1.2 update-source loopback 0
- E. R1ip route 10.1.1.2 255.255.255.255 192.168.1.2! router bgp 100neighbor 10.1.1.2 disable-connected-checkneighbor 10.1.1.2 update-source Loopback0

ANSWER: B E

Explanation:

The configurations containing `disable-connected-check`, a host route to the remote loopback, and `update-source loopback0` on each router establish the eBGP session. R1 and R2 use loopback addresses as their BGP neighbor addresses, so each router must first have IP reachability to the other router's loopback. The static /32 routes provide that reachability through the directly connected 192.168.1.0 network. The `update-source loopback0` command ensures that BGP packets originate from the local loopback address, matching the address configured by the remote peer.

Because these are eBGP peers, Cisco IOS normally performs a connected-neighbor check that expects the configured neighbor address to be directly connected. `disable-connected-check` removes that address-based restriction while retaining the normal eBGP TTL behavior. Since the loopbacks are reachable across the single directly connected transit link, this is sufficient without configuring eBGP multihop. Applying the corresponding configuration on both sides gives R1 and R2 reciprocal reachability, correct source addresses, and compatible eBGP session behavior. Cisco documents this command in the [Cisco IOS BGP Command Reference](#); see also Cisco's [BGP TTL Security documentation](#) for related eBGP TTL behavior.

QUESTION NO: 65

What are two characteristics of IPv6 Source Guard? (Choose two.)

- A. requires IPv6 snooping on Layer 2 access or trunk ports
- B. used in service provider deployments to protect DDoS attacks
- C. requires the user to configure a static binding
- D. requires that validate prefix be enabled
- E. recovers missing binding table entries

ANSWER: A E

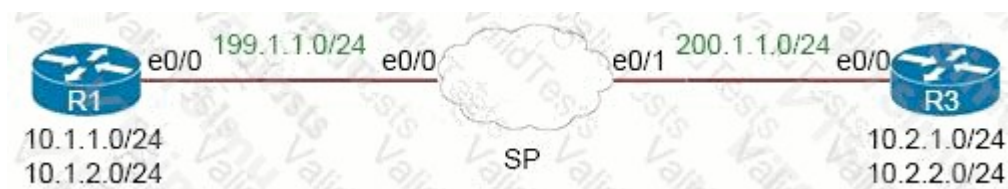
Explanation:

IPv6 Source Guard requires IPv6 snooping on Layer 2 access or trunk ports because it depends on valid IPv6-to-MAC binding information to determine whether traffic received on an interface is legitimate. DHCPv6 snooping supplies bindings learned from DHCPv6 exchanges, while related IPv6 First-Hop Security functions can contribute host information. The Source Guard policy is then attached to a Layer 2 interface and uses the binding data to filter packets whose source identity does not match an authorized host on that port. Cisco documents IPv6 Source Guard as an IPv6 First-Hop Security feature that uses source bindings to validate hosts and enforce source-address filtering.

IPv6 Source Guard also recovers missing binding table entries through IPv6 device tracking. Device tracking learns and maintains IPv6 neighbor and address information from IPv6 control-plane activity, allowing the switch to build or restore host binding information when an entry is not available from DHCPv6 snooping alone. This is particularly useful for hosts that use stateless addressing or whose DHCPv6-derived binding is absent. Cisco describes the interaction of Source Guard, snooping, and device tracking in its IPv6 First-Hop Security documentation: [Cisco IOS XE IPv6 First-Hop Security Configuration Guide](#) and [Cisco IOS XE IPv6 First-Hop Security Command Reference](#).

QUESTION NO: 66

Refer to the exhibit.



An engineer must configure a LAN-to-LAN IPsec VPN between R1 and the remote router. Which IPsec

Phase 1 configuration must the engineer use for the local router?

- A. `crypto isakmp policy 5 authentication pre-share encryption 3des hash sha group 2 ! crypto isakmp key cisco123 address 200.1.1.3`
- B. `crypto isakmp policy 5 authentication pre-share encryption 3des hash md5 group 2 ! crypto isakmp key cisco123 address 200.1.1.3`
- C. `crypto isakmp policy 5 authentication pre-share encryption 3des hash md5 group 2 ! crypto isakmp key cisco123 address 199.1.1.1`
- D. `crypto isakmp policy 5 authentication pre-share encryption 3des hash md5 group 2 ! crypto isakmp key cisco123 address 199.1.1.1`

ANSWER: A

Explanation:

The configuration using a pre-shared key for peer address 200.1.1.3, 3DES encryption, SHA hashing, and Diffie-Hellman group 2 is the required IKEv1 Phase 1 configuration for the local router. IKE Phase 1 establishes the ISAKMP security association that protects subsequent negotiation, so the local policy must use the same authentication method, encryption algorithm, hash algorithm, and DH group defined for the peer relationship in the exhibit. The `crypto isakmp key cisco123 address 200.1.1.3` command associates the pre-shared key with the remote VPN endpoint, rather than with the local router's own outside address. This lets R1 authenticate the peer when IKE negotiation begins. The policy number is locally significant; its role is to select the policy during negotiation, while the cryptographic parameters and peer key association determine interoperability. Cisco's IKEv1 configuration documentation describes defining ISAKMP policies and configuring pre-shared keys against the peer's IP address. See [Cisco IOS XE IKEv1 configuration guide](#) and [Cisco IKE command reference](#).

QUESTION NO: 67

LA

```
router ospf 1
 network 192.168.12.0 0.0.0.255 area 0
 network 172.16.1.0 0.0.0.255 area 0
```

NY

```
router ospf 1
 network 192.168.12.0 0.0.0.255 area 0
 network 172.16.2.0 0.0.0.255 area 0
!
interface E 0/0
 ip ospf authentication message-digest
 ip ospf message-digest-key 1 md5 Cisco123
```

Refer to the exhibit. The neighbor relationship is not coming up.

Which two configurations bring the adjacency up? (Choose two.)

A. LA

```
interface E 0/0
 ip ospf authentication-key Cisco123
```

B. NY

```
interface E 0/0 no ip ospf message-digest-key 1 md5 Cisco123 ip ospf authentication-key Cisco123
```

C. LA

```
interface E 0/0 ip ospf message-digest-key 1 md5 Cisco123
```

D. LA

```
router ospf 1
 area 0 authentication message-digest
```

E. NY

```
router ospf 1
 area 0 authentication message-digest
```

ANSWER: C D

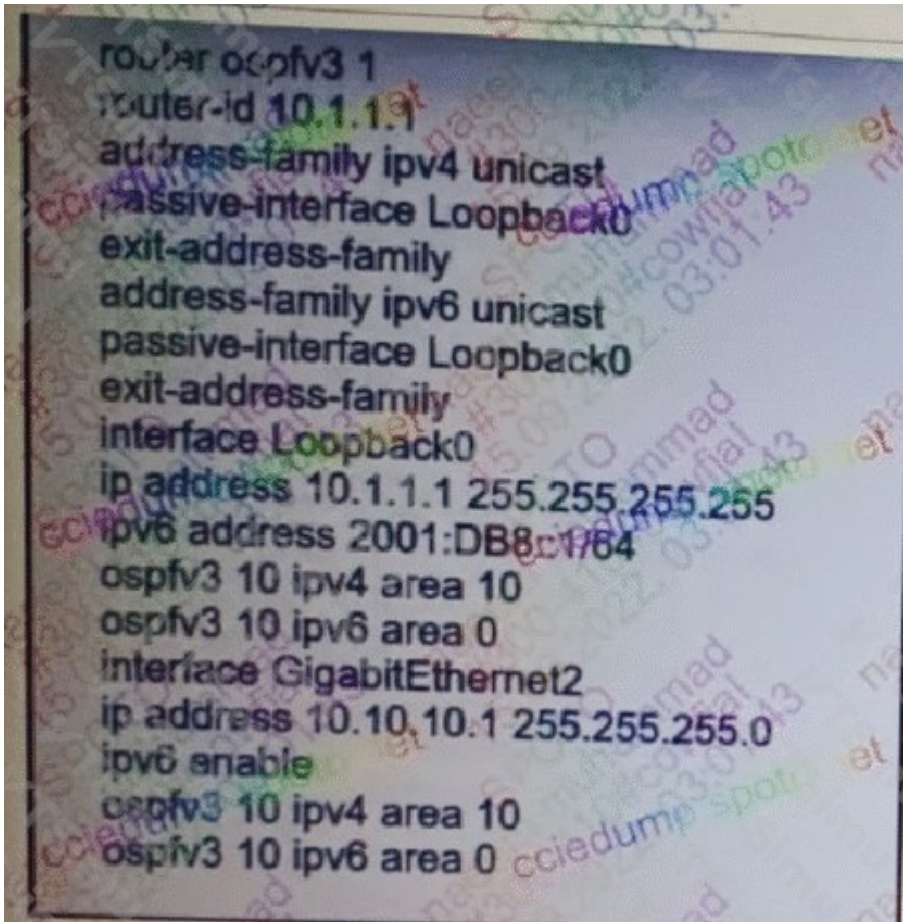
Explanation:

OSPF neighbors must use compatible authentication on the shared link before an adjacency can progress to the full state. The configuration `ip ospf message-digest-key 1 md5 Cisco123` under LA Ethernet0/0 supplies the MD5 key ID, authentication algorithm, and shared secret that LA uses to generate and validate OSPF message digests. The configuration `area 0 authentication message-digest` under LA's OSPF process enables message-digest authentication for interfaces in area 0, including the Ethernet0/0 link. Together, these commands make LA use MD5 authentication with key ID

1 and the password Cisco123, matching the MD5-based authentication required on the neighbor-facing segment shown in the exhibit. OSPF includes authentication information in its protocol packets, so both the authentication method and key material must align between directly connected neighbors. Once LA is configured for area-wide MD5 authentication and has the matching interface key, it can validate the neighbor's OSPF packets and form the adjacency. Cisco documents MD5 area authentication and interface message-digest key configuration in its [OSPF configuration guide](#) and [OSPF command reference](#).

QUESTION NO: 68

Refer to the exhibit.



```
router ospfv3 1
router-id 10.1.1.1
address-family ipv4 unicast
passive-interface Loopback0
exit-address-family
address-family ipv6 unicast
passive-interface Loopback0
exit-address-family
interface Loopback0
ip address 10.1.1.1 255.255.255.255
ipv6 address 2001:DB8::1 64
ospfv3 10 ipv4 area 10
ospfv3 10 ipv6 area 0
interface GigabitEthernet2
ip address 10.10.10.1 255.255.255.0
ipv6 enable
ospfv3 10 ipv4 area 10
ospfv3 10 ipv6 area 0
```

An administrator must configure the router with OSPF for IPv4 and IPv6 networks under a single process. The OSPF adjacencies are not established and did not meet the requirement. Which action resolves the issue?

- A. Replace OSPF process 10 on the interface with OSPF process 1, and configure an additional router ID with IPv6 address.
- B. Replace OSPF process 10 on the interface with OSPF process 1, for the IPv6 address and remove process router ID with IPv6 address.
- C. Replace OSPF process 10 on the interface with OSPF process 1, and remove process 10 from the global configuration.
- D. Replace OSPF process 10 on the interface with OSPF process 1 for the IPv4 address, and remove process 10 from the global configuration.

ANSWER: C

Explanation:

Replace OSPF process 10 on the interface with OSPF process 1, and remove process 10 from the global configuration. This makes the interface-level OSPFv3 configuration refer to the same OSPFv3 process that is configured globally. OSPFv3 supports IPv4 and IPv6 address families within one OSPFv3 routing process, so a single process can provide both protocol families when the corresponding address-family configuration is enabled. The process identifier applied under the interface

must match the existing global `router ospfv3 1` process; otherwise, the interface is associated with a different OSPFv3 process and cannot form the intended adjacency through the configured process. After changing the interface association to process 1 and removing the unused process 10 configuration, the router has one coherent OSPFv3 process for both IPv4 and IPv6, satisfying the stated requirement. Cisco's OSPFv3 documentation describes OSPFv3 support for IPv4 and IPv6 address families and the association of OSPFv3 with interfaces: [Cisco IOS XE OSPFv3 Configuration Guide](#). Cisco also documents the `ospfv3` interface command behavior in its command reference: [Cisco IOS IP Routing: OSPF Command Reference](#).

QUESTION NO: 69

The summary route is not shown in the RouterB routing table after this below configuration on Router_A

```
interface ethernet 0
description location ID:S4289T9E09F39
ip address 192.168.3.1 255.255.255.0
ip summary-address eigrp 1 172.16.80.0 255.255.240.0
```

Which Router_A configuration resolves the issue by advertising the summary route to Router B?

- ☐ interface loopback 0
ip address 172.16.96.1 255.255.255.0
interface Ethernet 0
ip address 192.168.3.1 255.255.255.0
ip summary-address eigrp 1 172.16.80.0 255.255.240.0
- ☐ interface loopback 0
ip address 172.16.81.1 255.255.255.0
interface Ethernet 0
ip address 192.168.3.1 255.255.255.0
ip summary-address eigrp 1 172.16.80.0 255.255.240.0
- ☒ interface loopback 0
ip address 172.16.79.1 255.255.255.0
interface Ethernet 0
ip address 192.168.3.1 255.255.255.0
ip summary-address eigrp 1 172.16.80.0 255.255.240.0
- ☐ interface loopback 0
ip address 172.18.81.1 255.255.255.0
interface Ethernet 0
ip address 192.168.3.1 255.255.255.0
ip summary-address eigrp 1 172.16.80.0 255.255.240.0

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

ANSWER: B

Explanation:

The configuration represented by "Option B" is the correct selection because it applies the required route-summary advertisement at Router_A toward Router_B. Route summarization is an outbound advertisement action: Router_A must originate the aggregate in the routing process and on the appropriate neighbor-facing interface or protocol boundary so that Router_B can receive and install the summary prefix. Once the aggregate is advertised, Router_B learns one less-specific route that represents the component networks covered by the summary, reducing its routing-table size and limiting routing-update scope. Cisco routing implementations also commonly install a local discard route for a configured summary, which prevents forwarding loops when traffic matches the aggregate but no more-specific route exists. The summary must be

configured with the intended network prefix and mask, and it must be associated with the routing protocol adjacency through which Router_B is reached. This ensures that the aggregate is included in updates sent from Router_A rather than remaining only a local routing-table concept. Cisco documents EIGRP manual summarization and its effects on route advertisement at [Cisco: EIGRP Route Summarization](#). General Cisco routing configuration guidance is available in the [Cisco IOS XE EIGRP Route Summary documentation](#).

QUESTION NO: 70

An engineer is trying to copy an IOS file from one router to another router by using TFTP.

Which two actions are needed to allow the file to copy? (Choose two.)

- A. Copy the file to the destination router with the `copy tftp: flash:` command
- B. Enable the TFTP server on the source router with the `tftp-server flash:filename` command
- C. TFTP is not supported in recent IOS versions, so an alternative method must be used
- D. Configure a user on the source router with the `username tftp password tftp` command
- E. Configure the TFTP authentication on the source router with the `tftp-server authentication local` command

ANSWER: A B

Explanation:

To transfer an IOS image directly between routers with TFTP, the router that currently holds the image must act as a TFTP server. Configuring `tftp-server flash:filename` publishes the specified file from local Flash so that a TFTP client can request it. The source router must also have the image present in Flash under the filename used in that command. Cisco IOS supports this built-in, lightweight TFTP-server capability for serving selected files; it does not require creation of a local username or TFTP authentication configuration for this transfer.

The receiving router acts as the TFTP client. On that router, `copy tftp: flash:` initiates the download, prompts for the remote server address and source filename, and writes the received image to local Flash. The routers must, of course, have IP reachability between the selected TFTP source and destination interfaces, and the destination must have sufficient available Flash storage for the image. After the copy completes, the image can be verified with commands such as `dir flash:` and an integrity check appropriate to the platform. Cisco documents the IOS file-copy workflow and the TFTP server configuration in its [Managing Cisco IOS Software Images](#) guide and [IP Application Services Command Reference](#).

QUESTION NO: 71

An engineer configured a leak-map command to summarize EIGRP routes and advertise specifically loopback 0 with an IP of 10.1.1.1.255.255.255.252 along with the summary route. After finishing configuration, the customer complained not receiving summary route with specific loopback address. Which two configurations will fix it? (Choose two.)

```
router eigrp 1
!
route-map Leak-Route deny 10
!
interface Serial 0/0
ip summary-address eigrp 1 10.0.0.0 255.0.0.0 leak-map Leak-Route
```

- A. Configure access-list 1 permit 10.1.1.0 0.0.0.3.
- B. Configure access-list 1 permit 10.1.1.1 0.0.0.252.
- C. Configure access-list 1 and match under route-map Leak-Route.
- D. Configure route-map Leak-Route permit 10 and match access-list 1.
- E. Configure route-map Leak-Route permit 20.

ANSWER: A D

Explanation:

Configuring `access-list 1 permit 10.1.1.0 0.0.0.3` correctly identifies the actual EIGRP prefix for the Loopback0 address. An interface address of 10.1.1.1 with a 255.255.255.252 mask belongs to the 10.1.1.0/30 network, and the wildcard mask 0.0.0.3 matches that complete /30 prefix. This is the route that must be selected for advertisement in addition to the EIGRP summary.

Configuring `route-map Leak-Route permit 10` and matching access list 1 makes the leak map explicitly permit the selected 10.1.1.0/30 route. EIGRP interface summarization normally advertises the summary while suppressing component routes covered by that summary. A leak map is applied to the summary configuration to exempt selected component prefixes from suppression. Therefore, the route map must have a permitting sequence that matches the access list selecting the connected Loopback0 network. With those commands in place, EIGRP advertises both the configured summary and the leaked 10.1.1.0/30 route. Cisco describes EIGRP summarization behavior in its [EIGRP troubleshooting documentation](#) and documents related EIGRP configuration commands in the [Cisco IOS XE EIGRP Configuration Guide](#).

QUESTION NO: 72

Refer to the exhibit. What does the `imp-null` tag represent in the MPLS VPN cloud?

```
Router# show tag-switching tdp bindings
(...)
tib entry: 10.10.10.1/32, rev 31
  local binding: tag: 18
  remote binding: tsr: 10.10.10.1:0, tag: imp-null
  remote binding: tsr: 10.10.10.2:0, tag: 18
  remote binding: tsr: 10.10.10.6:0, tag: 21
tib entry: 10.10.10.2/32, rev 22
  local binding: tag: 17
  remote binding: tsr: 10.10.10.2:0, tag: imp-null
  remote binding: tsr: 10.10.10.1:0, tag: 19
  remote binding: tsr: 10.10.10.6:0, tag: 22
```

- A. Pop the label
- B. Impose the label
- C. Include the EXP bit
- D. Exclude the EXP bit

ANSWER: A

Explanation:

Pop the label is correct. The `imp-null` tag denotes the MPLS implicit-null label, which has the reserved label value 3. It is not sent as an actual label in the labeled packet. Instead, when a downstream label-switching router advertises implicit null to its upstream neighbor, it is requesting penultimate-hop popping (PHP). The upstream, penultimate router removes the top MPLS transport label before forwarding the packet to the egress provider-edge router.

In an MPLS Layer 3 VPN, PHP avoids requiring the egress router to perform an unnecessary lookup and removal of the outer transport label. After the label is popped, the egress provider-edge router receives the packet with the VPN label still

available, when applicable, and uses that label to select the appropriate VPN routing and forwarding instance and forwarding context. Cisco documentation describes implicit null as a reserved MPLS label used to signal that the upstream router should pop the label, and documents PHP as the resulting forwarding behavior. See Cisco's [MPLS Layer 3 VPN configuration guide](#) and [MPLS label documentation](#).

QUESTION NO: 73

What are two benefits of configuring BGP route refresh capability? (Choose two.)

- A. It permits inbound BGP policy changes to be applied without resetting the session.
- B. It allows a router to request that its peer resend routing updates.
- C. It replaces the BGP OPEN message during neighbor establishment.
- D. It advertises routes that are suppressed by outbound policy.
- E. It automatically changes the local preference of refreshed routes.

ANSWER: A B

Explanation:

BGP route refresh permits a router to request that a peer resend its complete Adj-RIB-Out after inbound routing policy changes. The receiving router can then apply the revised policy to the refreshed updates without requiring the BGP session to be cleared. This avoids the routing disruption that can result from resetting an established BGP session.

Route refresh capability is negotiated between BGP peers through the BGP capability advertisement. When both peers support it, an administrator can use a soft inbound clear to request the updated advertisements dynamically. Route refresh does not eliminate the need for policy configuration, and it does not itself modify routes; it provides a method for reevaluating received updates. The capability is defined in [RFC 2918](#) and is described by Cisco in its [BGP Route Refresh and Soft Reconfiguration documentation](#).

QUESTION NO: 74

Refer to the exhibit.

```
NY
router ospf 1
 network 192.168.12.0 0.0.0.255 area 0
 network 172.16.2.0 0.0.0.255 area 0
!
interface E 0/0
 ip ospf authentication message-digest
 ip ospf message-digest-key 1 md5 Cisco123
```

The neighbor relationship is not coming up Which two configurations bring the adjacency up? (Choose two)

A. NY

router ospf 1

area 0 authentication message-digest

B. LA

```
interface E 0/0
```

```
ip ospf message-digest-key 1 md5 Cisco123
```

C. NY

```
interface E 0/0
```

```
no ip ospf message-digest-key 1 md5 Cisco123
```

```
ip ospf authentication-key Cisco123
```

D. LA

```
interface E 0/0
```

```
ip ospf authentication-key Cisco123
```

E. LA

```
router ospf 1
```

```
area 0 authentication message-digest
```

Answer: BE

The configuration on NY router is good for OSPF authentication. So we must enable OSPF

authentication on LA router with the following commands:

```
router ospf 1
```

```
area 0 authentication message-digest
```

```
interface E0/0
```

```
ip ospf message-digest-key 1 md5 Cisco123
```

A. LA

```
interface E 0/0
```

```
ip ospf message-digest-key 1 md5 Cisco123
```

B. NY

```
interface E 0/0
```

```
no ip ospf message-digest-key 1 md5 Cisco123
```

```
ip ospf authentication-key Cisco123
```

C. LA

```
interface E 0/0
```

```
ip ospf authentication-key Cisco123
```

D. LA

```
router ospf 1
```

```
area 0 authentication message-digest
```

ANSWER: A D

Explanation:

OSPF message-digest authentication is enabled consistently only when the area is configured to require message-digest authentication and each participating interface has a matching MD5 key. The NY router is already configured for MD5 authentication on the shared Ethernet segment, using key ID 1 and the password `Cisco123`. Therefore, LA must be

configured with `ip ospf message-digest-key 1 md5 Cisco123` under interface E0/0. This creates the matching per-interface MD5 key material needed to authenticate OSPF Hello packets and subsequent protocol packets.

LA must also have `area 0 authentication message-digest` under its OSPF process. This area-level command tells OSPF to use message-digest authentication for interfaces in area 0. With that setting and the matching E0/0 MD5 key, LA sends and accepts OSPF packets using the same authentication type, key ID, and key string as NY. The routers can then successfully validate each other's Hellos and progress through the OSPF neighbor-state machine to establish adjacency. Cisco's OSPF authentication guidance describes the requirement to configure MD5 authentication at the area level and define a message-digest key on the relevant interface: [Cisco OSPF Authentication Configuration](#).

QUESTION NO: 75

```
ip dhcp pool 1
network 200.30.30.0/24
default-router 200.30.30.100
lease 40
!
ip dhcp pool 2
network 200.30.40.0/24
default-router 200.30.40.100
lease 40
!
```

Refer to the exhibit. The server for the finance department is not reachable consistently on the 200.30.40.0/24 network and after every second month it gets a new IP address. What two actions must be taken to resolve this issue? (Choose two.)

- A. Configure the server to use DHCP on the network with default gateway 200.30.40.100.
- B. Configure the server with a static IP address and default gateway.
- C. Configure the router to exclude a server IP address.
- D. Configure the server to use DHCP on the network with default gateway 200.30.30.100.
- E. Configure the router to exclude a server IP address and default gateway.

ANSWER: B C

Explanation:

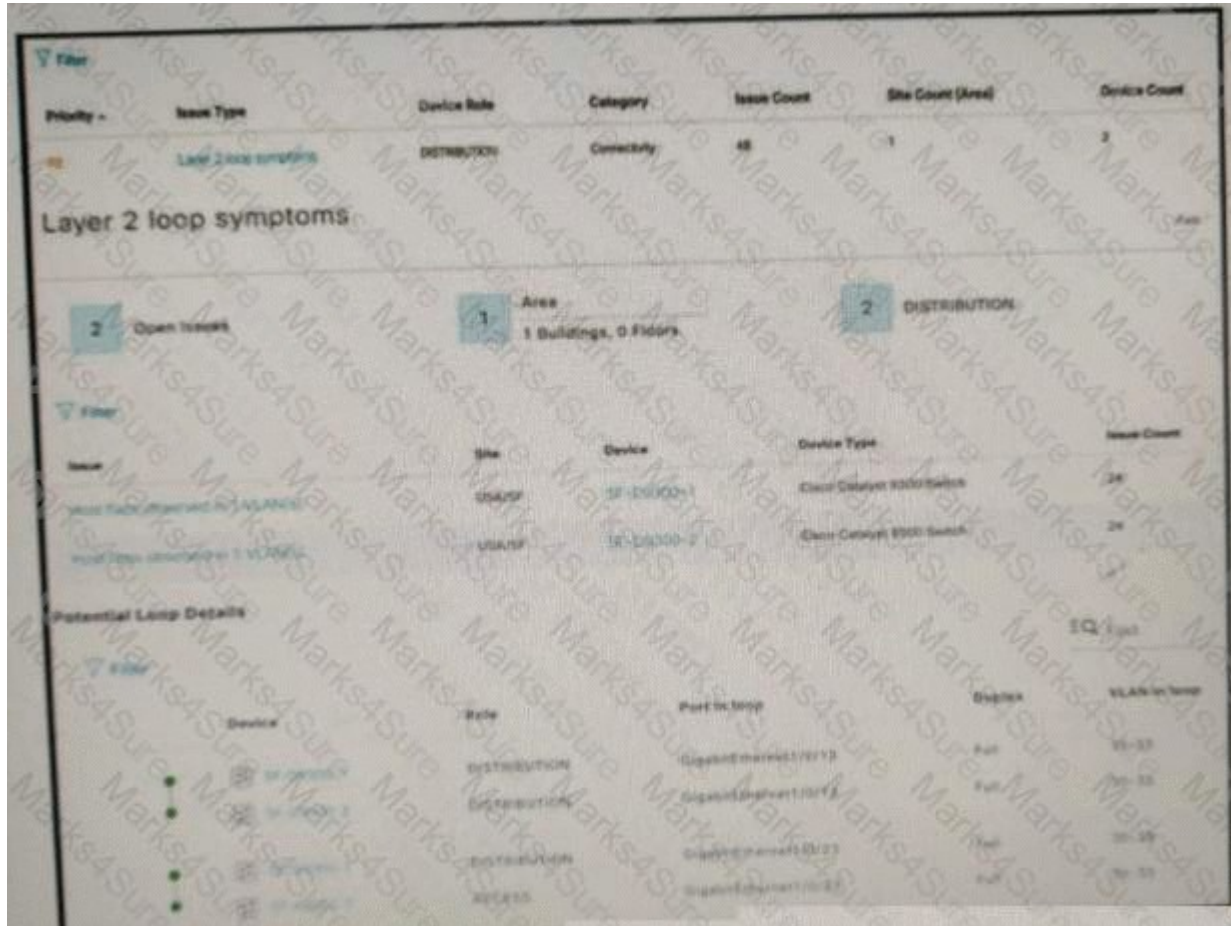
The server must be configured with a static IPv4 address in the 200.30.40.0/24 subnet and with the appropriate default gateway so that clients and services can always use the same destination address. A server is normally a fixed infrastructure endpoint, so a stable address is necessary for reliable access, DNS records, application configurations, and any access controls that reference its IP address. The observed two-month address change is consistent with the server receiving an address through DHCP and being subject to the DHCP lease process.

After selecting the server's static address, that address must be excluded from the router's DHCP allocation range. Cisco IOS DHCP exclusions reserve specified addresses so that the DHCP server never leases them to dynamic clients. This

prevents an address conflict in which another host receives the address assigned manually to the finance server. Together, static addressing for the server and a DHCP exclusion for that server address provide persistent reachability while allowing DHCP to continue serving the remaining client addresses on the subnet. Cisco documents DHCP address exclusion with the `ip dhcp excluded-address` command in its DHCP configuration guidance: [Cisco DHCP Configuration Examples and Troubleshooting](#) and [Cisco IOS XE DHCP Configuration Guide](#).

QUESTION NO: 76

Refer to the exhibit.



```
interface GigabitEthernet1/0/13
  switchport trunk allowed vlan 30-33
  switchport mode trunk
!
interface GigabitEthernet1/0/23
  switchport trunk allowed vlan 30-33
  switchport mode trunk
```

An engineer identifies a Layer 2 loop using DNAC. Which command fixes the problem in the SF-D9300-1 switch?

- A. no spanning-tree uplinkfast
- B. spanning-tree loopguard default
- C. spanning-tree backbonesfast
- D. spanning-tree portfast bpduguard default

ANSWER: D

Explanation:

`spanning-tree portfast bpduguard default` is the appropriate global configuration command to protect PortFast-enabled access ports from an unintended Layer 2 connection. PortFast is intended for edge devices, such as PCs, printers, and IP phones, which should not send BPDUs. If a switch, bridge, or accidental cable loop is connected to such a port, it can introduce BPDUs and create a switching loop risk. Enabling BPDU Guard by default causes the switch to place a PortFast-enabled port into the err-disabled state immediately when it receives a BPDU. This isolates the offending connection and stops the loop from continuing to affect the Layer 2 domain. The configuration is applied globally and automatically protects interfaces on which PortFast is enabled, avoiding the need to configure BPDU Guard individually on every access interface. Cisco recommends using BPDU Guard with PortFast to prevent topology loops caused by accidental switch-to-switch connections at the network edge. For command behavior and configuration guidance, see the [Cisco IOS XE Spanning Tree documentation](#) and Cisco's [BPDU Guard enhancement guidance](#).

QUESTION NO: 77

What are two characteristics of EIGRP stub routing? (Choose two.)

- A. It advertises connected and summary routes by default.
- B. It prevents neighboring routers from sending queries to the stub router.
- C. It must be configured on every router in the EIGRP autonomous system.
- D. It advertises all learned EIGRP routes by default.
- E. It converts all EIGRP routes into static routes.

ANSWER: A B

Explanation:

An EIGRP stub router advertises only limited reachability information to its neighbors. By default, it advertises connected and summary routes. This behavior prevents the stub router from being considered as a transit path and reduces EIGRP query scope during route loss or topology changes.

When a router is configured with `eigrp stub`, neighboring EIGRP routers do not send it queries for routes that the stub cannot provide. This reduces processing and bandwidth requirements at remote or hub-and-spoke locations. The stub router can be configured to advertise additional route types, such as static or redistributed routes, when required. Cisco documents the supported EIGRP stub route categories and query behavior in the [Cisco IOS XE EIGRP Stub Routing Configuration Guide](#).

QUESTION NO: 78 - (DRAG DROP)

Drag and drop the MPLS concepts from the left onto the descriptions on the right.

label edge router	allows an LSR to remove the label before forwarding the packet
label switch router	accepts unlabeled packets and imposes labels
forwarding equivalence class	group of packets that are forwarded in the same manner
penultimate hop popping	receives labeled packets and swaps labels

ANSWER:

label edge router	penultimate hop popping
label switch router	label edge router
forwarding equivalence class	forwarding equivalence class
penultimate hop popping	label switch router

Explanation:

MPLS forwarding is built around assigning packets to a forwarding equivalence class, or FEC, and then applying label operations as packets traverse the MPLS domain. A forwarding equivalence class is a group of packets that are handled in the same way by the MPLS network. Packets in the same FEC receive the same forwarding treatment, such as being sent over the same next hop and label-switched path. This classification lets MPLS forward traffic efficiently using labels rather than performing a full IP route lookup at every transit hop.

A label edge router operates at the boundary of the MPLS network. At ingress, it accepts an unlabeled IP packet, determines its FEC, and imposes an MPLS label stack so that the packet can enter the MPLS label-switched path. Within the MPLS core, a label switch router receives packets that already carry MPLS labels. It uses its label forwarding information to replace the incoming label with the appropriate outgoing label and forwards the packet toward its next hop.

Penultimate hop popping is the MPLS behavior in which the router immediately before the egress removes the top MPLS label before forwarding the packet. This means the egress router receives an unlabeled packet, avoiding an additional label lookup at the network edge. These functions together describe the normal MPLS lifecycle: classify traffic into an FEC, impose a label at ingress, swap labels through the core, and remove the label before egress. Cisco's MPLS overview discusses these label operations and router roles in its [MPLS configuration guide](#).

QUESTION NO: 79

The network administrator is tasked to configure R1 to authenticate telnet connections based on Cisco ISE using RADIUS. ISE has been configured with an IP address of 192.168.1.5 and with a network device pointing towards R1 (192.168.1.1) with a shared secret password of Cisco123. If ISE is down, the administrator should be able to connect using the local database with a username and password combination of admin/cisco123.

The administrator has configured the following on R1:

```
aaa new-model
!
username admin password cisco123
!
radius server ISE1
address ipv4 192.168.1.5
key Cisco123
!
aaa group server tacacs+ RAD-SERV
server name ISE1
!
aaa authentication login RAD-LOCAL group RAD-SERV
```

ISE has gone down. The Network Administrator is not able to Telnet to R1 when ISE went down. Which two configuration changes will fix the issue? (Choose two.)

- ☐ line vty 0 4
login authentication RAD-LOCAL
- ☐ line vty 0 4
login authentication default
- ☐ line vty 0 4
login authentication RAD-SERV
- ☐ aaa authentication login RAD-SERV group RAD-LOCAL local
- ☐ aaa authentication login RAD-LOCAL group RAD-SERV local

- A. Option A
- B. Option B
- C. aaa authentication login default group ISE local
- D. Option D
- E. login authentication default

ANSWER: C E

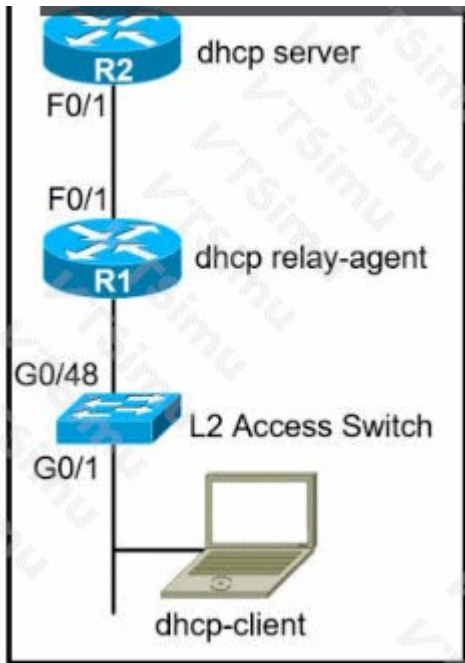
Explanation:

R1 requires an AAA login method list that uses the RADIUS server group as its primary authentication method and the local username database as its fallback method. The command `aaa authentication login default group ISE local` creates the default login method list with that ordered method sequence. When the RADIUS server is unavailable and cannot respond, IOS can proceed to the configured `local` method, where the locally configured `admin` account is evaluated. The local username and password must therefore match the required credentials, including password case.

The AAA method list does not affect Telnet merely by being defined globally; it must also be associated with the VTY lines that receive Telnet sessions. Applying `login authentication default` under the VTY configuration binds incoming remote-login attempts to the default AAA authentication list. Together, these commands ensure that Telnet first attempts Cisco ISE through the configured RADIUS server group and uses the router's local database when ISE is unreachable. Cisco documents AAA login method-list configuration and applying it to terminal lines in its [IOS XE AAA Configuration Guide](#) and the [AAA command reference](#).

QUESTION NO: 80

Refer to the exhibit.



The network administrator can see the DHCP discovery packet in R1. but R2 is not replying to the DHCP request. The R1 related interface is configured with the DHCP helper address. If the PC is directly connected to the FaO/1 interface on R2, the DHCP server assigns as IP address from the DHCP pool to the PC. Which two commands resolve this issue? (Choose two.)

- A. `service dhcp-relay` command on R1
- B. `ip dhcp option 82` command on R2
- C. `service dhcp` command on R1
- D. `ip dhcp relay information enable` command on R1
- E. `ip dhcp relay information trust-all` command on R2

ANSWER: C E

Explanation:

The `service dhcp` command on R1 enables the Cisco IOS DHCP service globally, including the DHCP relay-agent function used by an interface configured with `ip helper-address`. A relay agent receives the client broadcast, converts it to a unicast DHCP message toward the configured server, and forwards the server response back to the client subnet. Therefore, DHCP service must be enabled on the relay router for the helper-address-based relay operation to function.

The `ip dhcp relay information trust-all` command on R2 permits the DHCP server to accept relayed DHCP packets that contain relay-agent information (DHCP Option 82), regardless of the incoming interface's trust state. When R1 forwards the discovery with relay information, R2 must accept that relay-supplied information before it can select a pool and return a DHCP offer. Directly attached clients work because their DHCP broadcasts reach R2 locally and do not require relay-agent processing or acceptance of relay-added information. Together, enabling DHCP service on the relay and trusting

relay information on the server restores end-to-end DHCP allocation for the remote client. Cisco documents DHCP relay operation and the global DHCP service setting in the [Cisco IOS DHCP Configuration Guide](#); relay-information trust behavior is documented in the [Cisco IOS DHCP Command Reference](#).

QUESTION NO: 81

What are the two goals of micro BFD sessions? (Choose two.)

- A. The high bandwidth member link of a link aggregation group must run BFD
- B. Run the BFD session with 3x3 ms hello timer
- C. Continuity for each member link of a link aggregation group must be verified
- D. Eny member link on a link aggregation group must run BFD
- E. Each member link of a link aggregation group must run BFD.

ANSWER: C E

Explanation:

Micro BFD extends Bidirectional Forwarding Detection to individual member interfaces within a link aggregation group rather than treating the aggregate as a single forwarding path. Its purpose is to verify continuity independently for every member link, so a failure affecting one physical member can be detected even while the overall aggregated interface remains operational through its other members. This provides fast, granular failure detection for the paths that actually carry traffic in the bundle.

To achieve that per-member continuity verification, each member link of the link aggregation group must run a BFD session. Separate micro BFD sessions allow the devices to identify the specific failed member and remove or avoid that member while preserving traffic forwarding over healthy links. This behavior improves resiliency and prevents a physical-link fault from remaining undetected merely because the port channel itself is still up. Cisco describes BFD as a rapid, protocol-independent forwarding-path failure-detection mechanism in its [BFD technical documentation](#); the per-member LAG applicability is also standardized in [RFC 7130](#).

QUESTION NO: 82

Spoke# **show dmvpn**

Tunnel0, Type:Spoke, NHRP Peers:2,

Ent Peer NBMA Addr Peer Tunnel Add State UpDn Tm Attrb

1 172.18.16.2 192.168.1.1 UP 01:05:35 S

1 172.18.46.2 192.168.1.4 UP 00:00:25 D

Refer to the exhibit. An engineer has configured DMVPN on a spoke router.

What is the WAN IP address of another spoke router within the DMVPN network?

- A. 172.18.46.2
- B. 172.18.16.2
- C. 192.168.1.1

ANSWER: D**Explanation:**

192.168.1.4 is the WAN, or NBMA (Non-Broadcast Multi-Access), address for the remote spoke. DMVPN uses NHRP to maintain a mapping between a peer's tunnel address, which is used in the overlay, and its NBMA address, which is the reachable underlay/WAN address used to build the GRE/IPsec path. In the exhibit's NHRP information, the remote spoke is identified by its tunnel-side address while its associated NBMA address is shown as 192.168.1.4. Therefore, traffic sent directly to that spoke across the WAN must use 192.168.1.4 as the outer destination address. This distinction is fundamental to DMVPN operation: the tunnel address identifies the logical DMVPN peer, whereas the NBMA address identifies the physical endpoint reachable through the transport network. The hub and spokes use NHRP registrations and resolutions to learn these mappings dynamically, enabling spoke-to-spoke tunnels when required. Cisco's DMVPN overview describes NHRP's role in discovering the NBMA addresses of remote tunnel endpoints: [Cisco DMVPN configuration and troubleshooting guide](#). Cisco's NHRP documentation also explains the tunnel-to-NBMA address mapping used by DMVPN: [Cisco IOS XE NHRP and DMVPN documentation](#).

QUESTION NO: 83

Refer to the exhibit.

The network engineer configured the summarization of the RIP routes into the OSPF domain on R5 but still sees four different 172.16.0.0/24 networks on R4. Which action resolves the issue?

- A. R5(config)#router ospf 1
R5(config-router)#summary-address 172.16.0.0 255.255.252.0
- B. R4(config)#router ospf 99R4(config-router)#network 172.16.0.0 0.255.255.255 area 56R4(config-router)#area 56 range 172.16.0.0 255.255.255.0
- C. R4(config)#router ospf 1R4(config-router)#no areaR4(config-router)#summary-address 172.16.0.0 255.255.252.0
- D. R5(config)#router ospf 99R5(config-router)#network 172.16.0.0 0.255.255.255 area 56R5(config-router)#area 56 range 172.16.0.0 255.255.255.0

ANSWER: A**Explanation:**

R5(config)#router ospf 1
R5(config-router)#summary-address 172.16.0.0 255.255.252.0 is correct because R5 is the autonomous system boundary router that redistributes RIP routes into OSPF. Routes injected through redistribution are OSPF external routes, and external-route aggregation is configured on the ASBR with the OSPF `summary-address` command. The mask 255.255.252.0 represents a /22 prefix, which aggregates the four consecutive /24 networks beginning at 172.16.0.0: 172.16.0.0/24 through 172.16.3.0/24. R5 then advertises the resulting 172.16.0.0/22 external summary into OSPF, allowing R4 to learn one summarized prefix rather than four individual redistributed RIP prefixes. This behavior applies even when the ASBR is connected to a not-so-stubby area; the relevant distinction is that redistribution occurs at R5, so the external summary must also be configured there. Cisco documents that OSPF route redistribution and external-route summarization are performed by an ASBR using `summary-address`. See Cisco's [OSPF route redistribution documentation](#) and the [Cisco OSPF command reference](#).

QUESTION NO: 84

Which Ipv6 first-hop security feature helps to minimize denial of service attacks?

- A. IPv6 Router Advertisement Guard
- B. IPv6 Destination Guard

C. DHCPv6 Guard

D. IPv6 MAC address filtering

ANSWER: B

Explanation:

IPv6 Destination Guard helps minimize denial-of-service attacks by validating IPv6 destination addresses at the first-hop switch. It uses the IPv6 snooping binding database, which can contain address and port information learned through mechanisms such as DHCPv6 snooping and Neighbor Discovery snooping, to identify legitimate destinations on an access segment. The switch can then limit forwarding toward destinations that do not have a valid binding. This prevents traffic from being indiscriminately sent to unused, unknown, or unauthorized IPv6 addresses on protected access ports.

By restricting traffic to validated destination addresses, IPv6 Destination Guard reduces the opportunity for an attacker to generate large volumes of traffic toward arbitrary addresses in an attempt to consume endpoint, link, or switch resources. This is particularly useful at the access layer, where enforcing the policy close to the traffic source prevents unwanted packets from reaching hosts. Cisco documents IPv6 Destination Guard as part of its IPv6 First-Hop Security feature set and describes its use of IPv6 snooping information to protect hosts and mitigate traffic-based attacks. See the [Cisco IOS XE IPv6 First-Hop Security Configuration Guide](#) and [Cisco IPv6 First-Hop Security overview](#).

QUESTION NO: 85

Refer to the exhibit.



The AP status from Cisco DNA Center Assurance Dashboard shows some physical connectivity issues from access switch interface G1/0/14. Which command generates the diagnostic data to resolve the physical connectivity issues?

- A. test cable-diagnostics tdr interface GigabitEthernet1/0/14
- B. Check cable-diagnostics tdr interface GigabitEthernet1/0/14
- C. show cable-diagnostics tdr interface GigabitEthernet1/0/14
- D. Verify cable-diagnostics tdr interface GigabitEthernet1/0/14

ANSWER: A

Explanation:

test cable-diagnostics tdr interface GigabitEthernet1/0/14 is correct because it initiates a time-domain reflectometer (TDR) test on the affected switch port. TDR sends a signal through the attached copper Ethernet cable and

analyzes reflections caused by changes in impedance. This allows the switch to generate diagnostic information that can identify common Layer 1 cable faults, including an open circuit, a short circuit, or an impedance mismatch, and can estimate the distance from the port to the fault. Running the test against the interface identified by Cisco DNA Center provides the physical-layer evidence needed to investigate the reported connectivity issue. After the test has completed, its results can be reviewed with the related `show cable-diagnostics tdr interface GigabitEthernet1/0/14` command; however, the `test` command is the action that actually starts data collection and generates the diagnostic result. TDR support and behavior can vary by switch platform, interface type, and whether the port has an active link, so the applicable platform documentation should be consulted when interpreting results. Cisco documents cable diagnostics and TDR operation in its [cable-diagnostics technical note](#) and provides interface command guidance in the [Cisco IOS interface command reference](#).

QUESTION NO: 86

What is a function of IPv6 Source Guard?

- A. It works with address glean or ND to find existing addresses.
- B. It inspects ND and DHCP packets to build an address binding table.
- C. It denies traffic from known sources and allocated addresses.
- D. It notifies the ND protocol to inform hosts if the traffic is denied by it.

ANSWER: A

Explanation:

“It works with address glean or ND to find existing addresses.” is correct because IPv6 Source Guard enforces source-address validation on traffic entering an enabled Layer 2 interface. Its filtering decision is based on IPv6-to-interface bindings maintained in the device binding table. IPv6 Source Guard can use IPv6 address gleaning or IPv6 Neighbor Discovery inspection to learn valid IPv6 addresses present on the link and populate the information needed for enforcement. Once an address is learned and associated with the appropriate interface, Source Guard permits traffic that matches the valid binding. This provides protection against traffic using an unrecognized or spoofed IPv6 source address while allowing legitimately learned hosts to communicate. The feature is therefore an enforcement mechanism that relies on address-learning capabilities, rather than being the protocol process that independently performs all discovery itself. Cisco’s IPv6 First-Hop Security documentation describes IPv6 Source Guard as filtering IPv6 traffic according to source-address bindings and identifies IPv6 address gleaning and Neighbor Discovery inspection as mechanisms used with the feature. See the [Cisco IOS XE IPv6 First-Hop Security Configuration Guide](#) and the [Cisco IOS IPv6 command reference](#).

QUESTION NO: 87

Refer to the exhibits.

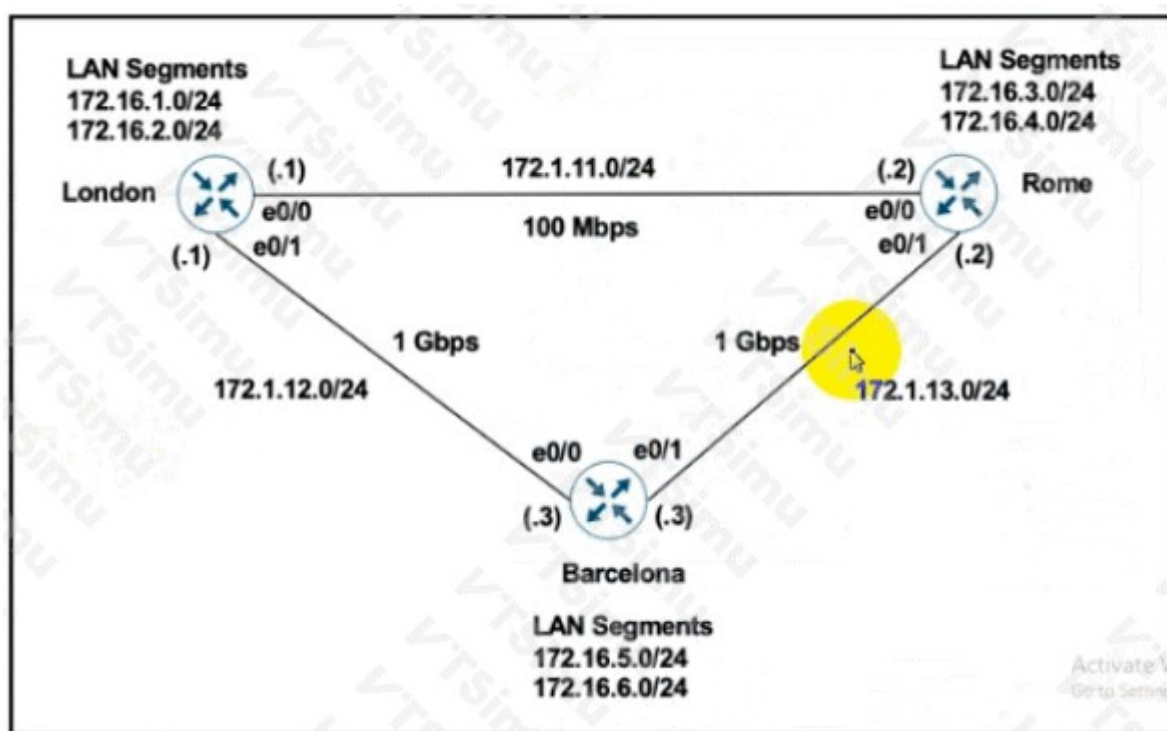
London - "show ip route" output

Gateway of last resort is not set

```
172.1.0.0/16 is variably subnetted, 5 subnets, 2 masks
C   172.1.11.0/24 is directly connected, Ethernet0/0
L   172.1.11.1/32 is directly connected, Ethernet0/0
C   172.1.12.0/24 is directly connected, Ethernet0/1
L   172.1.12.1/32 is directly connected, Ethernet0/1
D   172.1.13.0/24 [90/76800] via 172.1.11.2, 00:00:50, Ethernet0/0
172.16.0.0/16 is variably subnetted, 8 subnets, 2 masks
C   172.16.1.0/24 is directly connected, Loopback0
L   172.16.1.1/32 is directly connected, Ethernet0/0
C   172.16.2.0/24 is directly connected, Loopback1
L   172.16.2.1/32 is directly connected, Loopback1
R   172.16.3.0/24 [120/1] via 172.1.11.2, 00:00:08, Ethernet0/0
R   172.16.4.0/24 [120/1] via 172.1.11.2, 00:00:08, Ethernet0/0
D   172.16.5.0/24 [90/156160] via 172.1.12.3, 00:00:50, Ethernet0/1
D   172.16.6.0/24 [90/156160] via 172.1.12.3, 00:00:50, Ethernet0/1
```

Rome - "show run | section router" output

```
router eigrp 111
 network 172.1.0.0
 network 172.16.0.0
 no auto-summary
```



London must reach Rome using a faster path via EIGRP if all the links are up but it failed to take this path Which action resolves the issue?

- A. Increase the bandwidth of the link between London and Barcelona
- B. Use the network statement on London to inject the 172.16.X.0/24 networks into EIGRP.
- C. Change the administrative distance of RIP to 150
- D. Use the network statement on Rome to inject the 172.16.X.0/24 networks into EIGRP

ANSWER: D

Explanation:

Use the network statement on Rome to inject the 172.16.X.0/24 networks into EIGRP. For London to select an EIGRP route to a destination located behind Rome, Rome must advertise its directly connected destination networks into the EIGRP routing process. The EIGRP `network` command both enables EIGRP on interfaces matching the specified address range and causes connected networks on those interfaces to be advertised to EIGRP neighbors. Once Rome originates the relevant 172.16.X.0/24 prefixes, the EIGRP advertisements can travel across the London–Barcelona–Rome path. London can then install the EIGRP-learned route and use that path according to normal EIGRP route selection.

This change addresses route reachability and route advertisement at the source of the missing EIGRP prefix. It does not merely alter a path attribute: it ensures that the Rome-side destination networks exist in the EIGRP topology and routing tables so they can be propagated to London. Cisco documents that EIGRP network statements identify the interfaces and networks that participate in EIGRP routing. See Cisco's [EIGRP configuration guide](#) and the [Cisco EIGRP overview](#).

QUESTION NO: 88

Which feature drops packets if the source address is not found in the snooping table?

- A. IPv6 Source Guard
- B. IPv6 Destination Guard
- C. IPv6 Prefix Guard
- D. Binding Table Recovery

ANSWER: A

Explanation:

IPv6 Source Guard is correct because it protects a Layer 2 access port by validating the IPv6 source address of received traffic against the IPv6 First-Hop Security binding information learned through IPv6 snooping. The snooping binding table associates valid host details, such as the IPv6 address, MAC address, VLAN, and interface. When IPv6 Source Guard is enabled on an untrusted host-facing port, packets whose source address does not match an entry in that binding table are discarded. This prevents a device from spoofing an IPv6 source address or sending traffic before it has established a valid binding through the supported address-assignment and neighbor-discovery processes. IPv6 Source Guard can therefore be used as an access-layer anti-spoofing control and is typically deployed alongside IPv6 ND inspection and DHCPv6 guard as part of IPv6 First-Hop Security. Cisco documents that IPv6 Source Guard filters IPv6 traffic based on the IPv6 source address and the bindings maintained by the device. See Cisco's [IPv6 Source Guard configuration guide](#) and the [IPv6 First-Hop Security overview](#).

QUESTION NO: 89

An engineer configured two routers connected to two different service providers using BGP with default attributes. One of the links is presenting high delay, which causes slowness in the network. Which BGP attribute must the engineer configure to avoid using the high-delay ISP link if the second ISP link is up?

- A. LOCAL_PREF
- B. MED
- C. WEIGHT
- D. AS-PATH

ANSWER: A

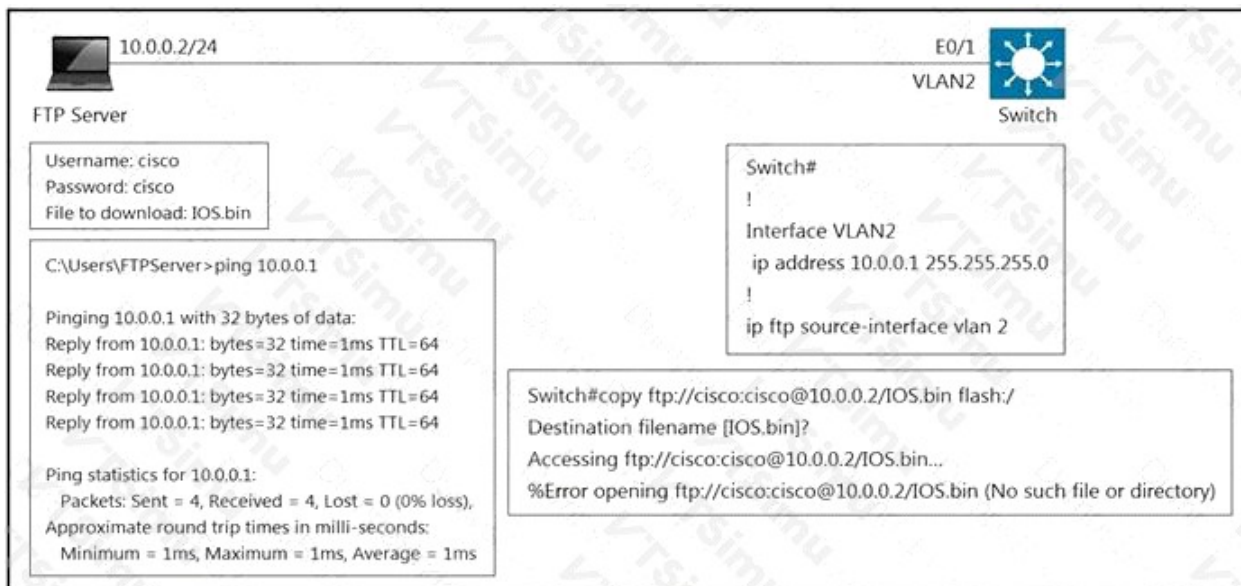
Explanation:

LOCAL_PREF is the appropriate attribute for selecting the preferred outbound ISP link throughout the enterprise autonomous system. The engineer can apply an inbound BGP policy on routes learned from each provider, assigning a higher LOCAL_PREF to routes received through the healthy, lower-delay ISP and a lower value to equivalent routes learned through the high-delay provider. BGP then prefers the path with the highest LOCAL_PREF before evaluating later path-selection criteria. As a result, traffic destined for external networks exits through the preferred ISP whenever that route is available; the alternate provider becomes the path used when the preferred link or its learned routes are unavailable.

LOCAL_PREF is carried to iBGP peers, which is important when the ISP connections terminate on different edge routers. This propagates the same exit-path preference to routers across the autonomous system and provides consistent outbound routing behavior rather than a decision that is limited to one router. Cisco documents that the default local preference is 100 and that higher values are preferred in BGP best-path selection. See Cisco's [BGP Best Path Selection Algorithm](#) and [BGP Path Selection and Route Manipulation](#).

QUESTION NO: 90

Refer to the exhibit.



An engineer cannot copy the IOS.bin file from the FTP server to the switch.

Which action resolves the issue?

- A. Allow file permissions to download the file from the FTP server.
- B. Add the IOS.bin file, which does not exist on FTP server.
- C. Make memory space on the switch flash or USB drive to download the file.
- D. Use the copy flash:/ ftp://cisco@10.0.0.2/IOS.bin command.

ANSWER: B

Explanation:

Add the IOS.bin file, which does not exist on FTP server. is the required action because the FTP copy operation identifies that the requested source filename cannot be found at the specified FTP location. A Cisco IOS copy operation requires the source URL, server reachability, credentials, and exact remote filename to resolve successfully before any transfer to local flash can begin. The filename is case-sensitive when the FTP server's underlying filesystem is case-sensitive, and it must also be present in the FTP user's accessible directory. Therefore, the engineer should upload the intended IOS image to the FTP server, or ensure that the source filename entered in the copy command exactly matches the image stored there. Once IOS.bin is present and accessible at that FTP path, the switch can retrieve it and write it to its selected local destination. Cisco documents file-copy operations and supported URL-based file systems in its IOS XE fundamentals documentation: [Cisco IOS XE Managing Files](#). Cisco's command reference also describes the `copy` command and its source and destination file-system behavior: [Cisco IOS XE Fundamentals Command Reference](#).

QUESTION NO: 91

Refer to the exhibit.

```
Spoke# show dmvpn
Tunnel0, Type:Spoke, NHRP Peers:2,
# Ent Peer NBMA Addr Peer Tunnel Add State UpDn Tm Attrb
-----
1 172.18.16.2 192.168.1.1 UP 01:05:35 S
1 172.18.46.2 192.168.1.4 UP 00:00:25 D
```

An engineer has configured DMVPN on a spoke router. What is the WAN IP address of another spoke router within the DMVPN network?

- A. 172.18.46.2
- B. 192.168.1.4
- C. 172.18.16.2
- D. 192.168.1.1

ANSWER: A

Explanation:

172.18.46.2 is the WAN, or NBMA (Non-Broadcast Multiaccess), address of the remote spoke. In a DMVPN deployment, each spoke has a tunnel interface address used for the overlay network and a separate physical/WAN address used as the NBMA endpoint in the underlay network. NHRP maintains the mapping between a remote spoke's overlay tunnel address and its NBMA address. The relevant DMVPN/NHRP output identifies the remote spoke's tunnel protocol address and associates it with 172.18.46.2 as the NBMA address. That WAN address is what the local spoke uses to build a direct dynamic GRE/IPsec spoke-to-spoke tunnel after NHRP resolution, avoiding unnecessary forwarding through the hub. Cisco describes NHRP as the protocol that allows a router to determine the NBMA address of a remote peer from its protocol address, which is the key mechanism enabling dynamic spoke-to-spoke connectivity in DMVPN. See Cisco's [DMVPN configuration guide](#) and its [DMVPN overview](#) for details on NHRP address mappings and NBMA peer resolution.

QUESTION NO: 92

Refer to the exhibit.

```
*Sep 26 19:50:43.504: SNMP: Packet received via UDP from
192.168.1.2 on GigabitEthernet0/1SrParseV3SnmpMessage: No
matching Engine ID.

SrParseV3SnmpMessage: Failed.
SrDoSnmp: authentication failure, Unknown Engine ID

*Sep 26 19:50:43.504: SNMP: Report, reqid 29548, errstat 0,
erridx 0
internet.6.3.15.1.1.4.0 = 3
*Sep 26 19:50:43.508: SNMP: Packet sent via UDP to 192.168.1.2
process_mgmt_req_int: UDP packet being de-queued
```

Which two commands provide the administrator with the information needed to resolve the issue? (Choose two.)

- A. snmp user
- B. debug snmp engine-id
- C. debug snmpv3 engine-id
- D. debug snmp packets
- E. showsnmpv3 user

ANSWER: B D

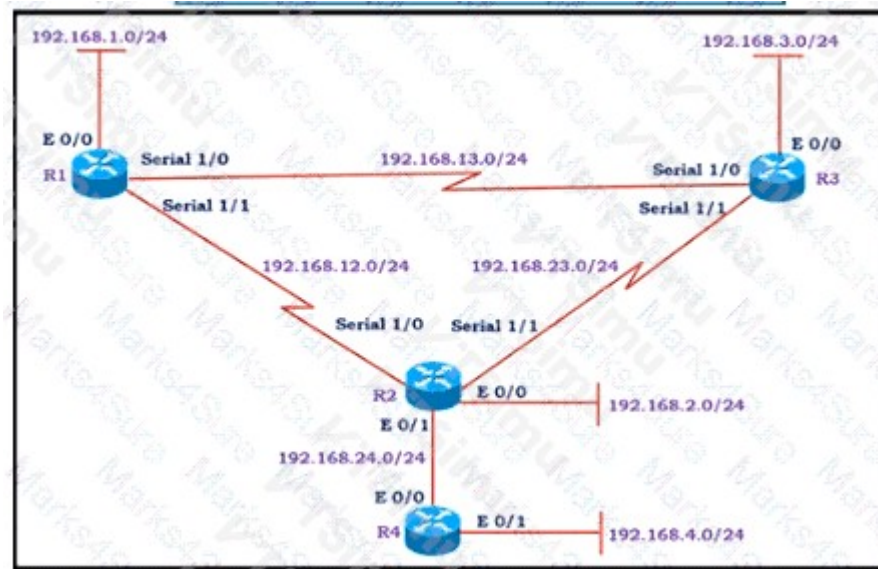
Explanation:

`debug snmp engine-id` provides real-time details about SNMPv3 engine-ID processing. The SNMPv3 engine ID is fundamental to user-based security because it identifies the authoritative SNMP engine and is used in localized authentication and privacy key processing. This debug is therefore useful when the displayed behavior indicates that an SNMPv3 peer is using, learning, or rejecting an engine ID unexpectedly.

`debug snmp packets` supplies packet-level diagnostic output for SNMP traffic. It lets the administrator observe the SNMP exchanges occurring between the manager and agent, including requests, responses, and SNMPv3 security-related processing. Used with engine-ID debugging, it provides the operational evidence needed to correlate a received SNMP packet with the engine-ID handling that occurs and identify the source of the communication problem. Debug commands should be enabled carefully on production devices because their output can be verbose and affect CPU utilization. Cisco documents SNMP debugging and SNMPv3 engine-ID behavior in its [SNMP Version 3 Configuration Guide](#) and the [Cisco IOS SNMP Command Reference](#).

QUESTION NO: 93

Refer to the exhibit.



Show IP route on R1

```
192.168.1.0/24 is variably subnetted, 2 subnets, 2 masks
C   192.168.1.0/24 is directly connected, Ethernet0/0
L   192.168.1.1/32 is directly connected, Ethernet0/0
D   192.168.2.0/24 [90/2297856] via 192.168.12.2, 00:02:14, Serial1/1
S   192.168.3.0/24 [1/0] via 192.168.12.2
192.168.12.0/24 is variably subnetted, 2 subnets, 2 masks
C   192.168.12.0/24 is directly connected, Serial1/1
L   192.168.12.1/32 is directly connected, Serial1/1
192.168.13.0/24 is variably subnetted, 2 subnets, 2 masks
C   192.168.13.0/24 is directly connected, Serial1/0
L   192.168.13.1/32 is directly connected, Serial1/0
D   192.168.23.0/24 [90/2681856] via 192.168.13.3, 00:06:38, Serial1/0
    [90/2681856] via 192.168.12.2, 00:06:38, Serial1/1
```

All the serial between R1, R2, and R3 have the Same bandwidth. User on the 192.168.1.0/24 network report slow response times while they access resource on network 192.168.3.0/24. When a traceroute is run on the path. It shows that the packet is getting forwarded via R2 to R3 although the link between R1 and R3 is still up. What must the network administrator to fix the slowness?

- A. Change the Administrative Distance of EIGRP to 5.
- B. Add a static route on R1 using the next hop of R3.
- C. Remove the static route on R1.
- D. Redistribute the R1 route to EIGRP

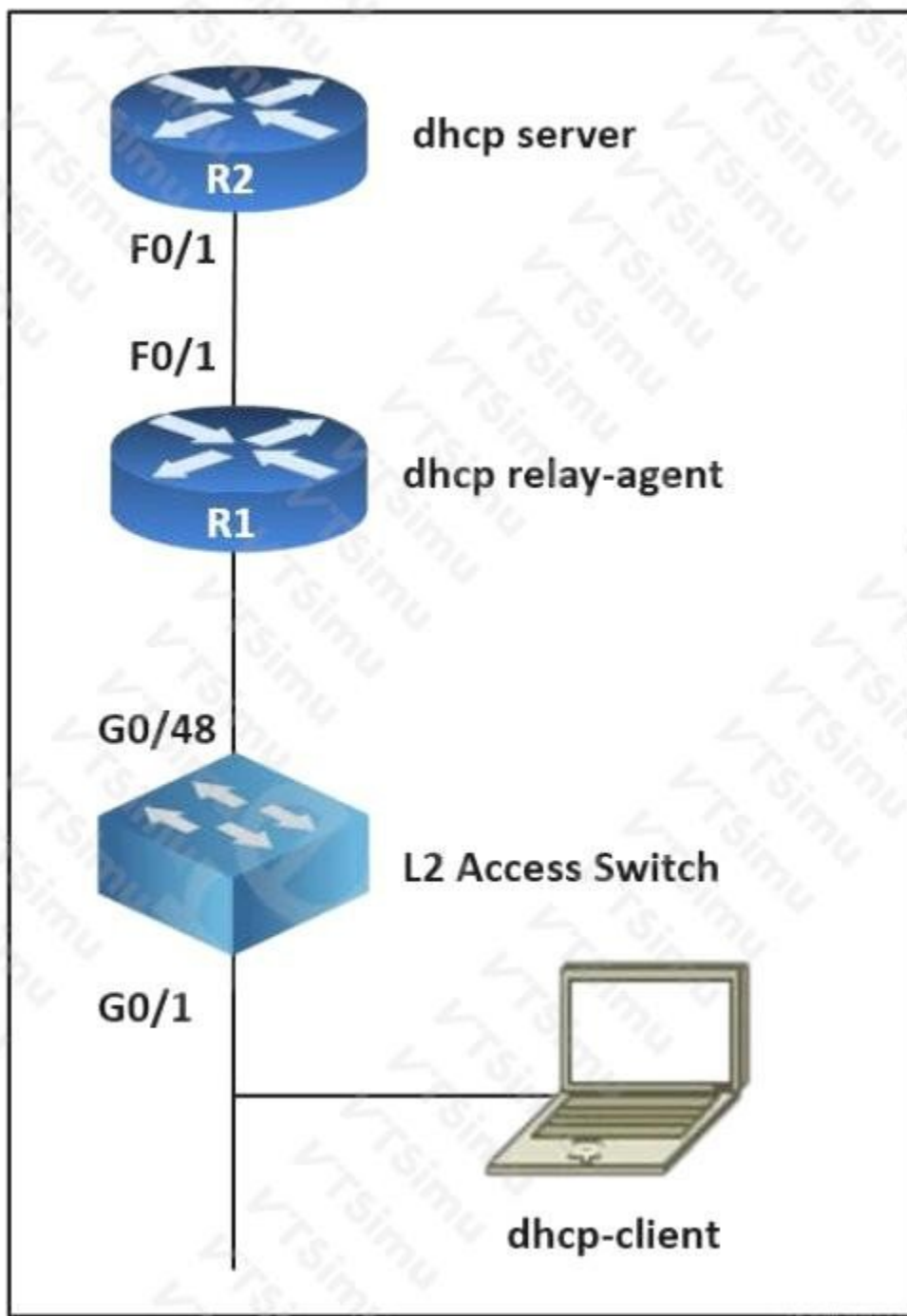
ANSWER: C

Explanation:

Remove the static route on R1. The static route is causing R1 to select a path through R2 for traffic destined for the remote 192.168.3.0/24 network, even though R1 has an available direct serial link to R3. By default, a static route has an administrative distance of 1, whereas an internal EIGRP route has an administrative distance of 90. Therefore, R1 installs the static route in its routing table ahead of the EIGRP-learned route, regardless of the EIGRP metric that would favor the direct R1-to-R3 connection. Removing the unnecessary static route permits R1 to use its EIGRP route selection again. Because the serial links have equal bandwidth, the direct R1-to-R3 path has a lower composite EIGRP metric than a route that must traverse both R1-to-R2 and R2-to-R3 links. Traffic will then use the direct link, reducing the number of hops and avoiding the unnecessary transit through R2. Cisco documents static-route configuration and behavior at [Cisco IOS IP Routing: Static Routes](#). Cisco's EIGRP overview, including route selection concepts, is available at [Cisco EIGRP Technical Documentation](#).

QUESTION NO: 94

Refer to the exhibit.



The network administrator can see the DHCP discovery packet in R1, but R2 is not replying to the DHCP request. The R1 related interface is configured with the DHCP helper address. If the PC is directly connected to the Fa0/1 interface on R2, the DHCP server assigns as IP address from the DHCP pool to the PC. Which two commands resolve this issue? (Choose two.)

- A. service dhcp-relay command on R1
- B. ip dhcp relay information enable command on R1
- C. ip dhcp option 82 command on R2
- D. service dhcp command on R1
- E. ip dhcp relay information trust-all command on R2

ANSWER: D E

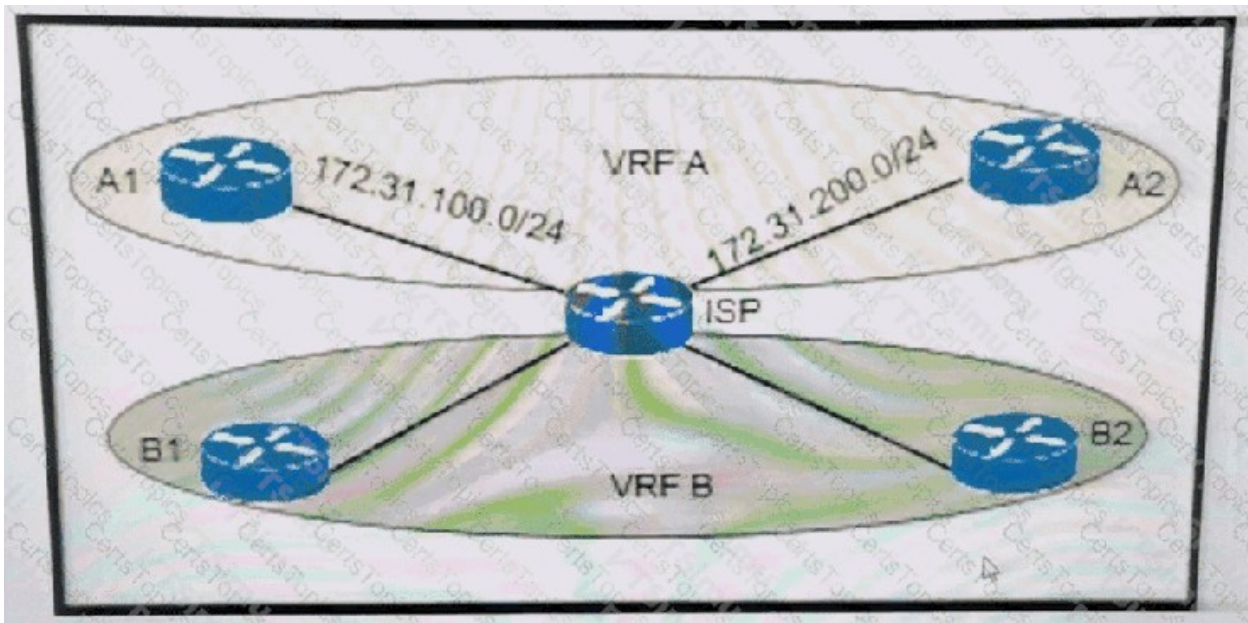
Explanation:

`service dhcp` on R1 enables the Cisco IOS DHCP service, including DHCP relay-agent operation. A router interface configured with `ip helper-address` must have the DHCP service enabled so that client broadcasts received on the client-facing interface are relayed as unicast DHCP messages toward the configured DHCP server. This permits R1 to forward the client request and populate the gateway address used by the server to select the appropriate address pool.

`ip dhcp relay information trust-all` on R2 permits its Cisco IOS DHCP server process to accept DHCP relay information, including DHCP Option 82, from relay agents. Option 82 identifies relay-agent and circuit information and can be present in a request relayed by R1. Trusting relay information allows R2 to process the relayed request and return the DHCPOFFER/DHCPACK through R1, while R2 can still allocate addresses successfully from the pool associated with the client subnet. Cisco documents DHCP relay and server operation in its [IOS XE DHCP Configuration Guide](#); DHCP relay-agent information is standardized in [RFC 3046](#).

QUESTION NO: 95

Refer to the exhibit. The ISP router is fully configured for customer A and customer B using the VRF-Lite feature. What is the minimum configuration required for customer A to communicate between routers A1 and A2?



A. A1
interface fa0/0
description To -> ISP
ip address 172.31.100.1 255.255.255.0
no shutdown
router ospf 100
network 172.31.100.1 0.0.0.255 area 0

A2
interface fa0/0
description To -> ISP
ip address 172.31.200.1 255.255.255.0
no shutdown
router ospf 100
network 172.31.200.1 0.0.0.255 area 0

B. A1
interface fa0/0
description To -> ISP
ip vrf forwarding A
ip address 172.31.100.1 255.255.255.0
no shutdown
router ospf 100
network 172.31.100.1 0.0.0.255 area 0

```
A2
interface fa0/0
description To -> ISP
ip vrf forwarding A
ip address 172.31.200.1 255.255.255.0
no shutdown
router ospf 100
network 172.31.200.1 0.0.0.255 area 0
```

```
C. A1
interface fa0/0
description To -> ISP
ip address 172.31.200.1 255.255.255.0
no shutdown
router ospf 100
network 172.31.200.1 0.0.0.255 area 0
```

```
A2
interface fa0/0
description To -> ISP
ip address 172.31.100.1 255.255.255.0
no shutdown
router ospf 100
network 172.31.100.1 0.0.0.255 area 0
```

```
D. A1
interface fa0/0
description To -> ISP
ip vrf forwarding A
ip address 172.31.100.1 255.255.255.0
no shutdown
router ospf 100 vrf A
network 172.31.200.1 0.0.0.255 area 0
```

```
A2
interface fa0/0
description To -> ISP
ip vrf forwarding A
ip address 172.31.100.1 255.255.255.0
no shutdown
router ospf 100 vrf A
network 172.31.200.1 0.0.0.255 area 0
```

ANSWER: A

Explanation:

A1: interface fa0/0; ip address 172.31.100.1 255.255.255.0; router ospf 100; network 172.31.100.1 0.0.0.255 area 0. A2: interface fa0/0; ip address 172.31.200.1 255.255.255.0; router ospf 100; network 172.31.200.1 0.0.0.255 area 0. is the minimum required configuration. The ISP-facing interfaces on A1 and A2 use the global routing table, because the ISP router—not either customer router—is the device that provides customer separation through VRF-Lite. A1 must address its link in the 172.31.100.0/24 subnet, while A2 must address its link in the 172.31.200.0/24 subnet, matching the respective Customer A-facing ISP interfaces. Enabling OSPF process 100 on each connected interface allows each customer router to establish an adjacency with the ISP's OSPF process that operates in VRF A. The ISP can then advertise routes between the two Customer A attachment networks within that VRF, while maintaining isolated routing and forwarding for the customer. A VRF configuration is unnecessary on A1 and A2 for this design because they each connect to only one customer context; their required participation is standard IP addressing and OSPF on the WAN-facing interface. Cisco describes VRF-Lite as maintaining separate routing and forwarding tables on the device providing virtualization: [Cisco IOS XE VRF Configuration Guide](#). OSPF forms neighbor relationships over enabled, matching IP interfaces as documented in [Cisco OSPF documentation](#).