



CISCO ENARSI 300-410 Exam Questions

Total Questions: 600+

Demo Questions: 35

Version: Updated for 2025

Prepared and Verified by Cert Empire – Your Trusted IT Certification Partner

**For Access to the full set of Updated Questions – Visit:
[CISCO ENARSI 300-410 Exam Questions](#) by Cert Empire**

Question: 1

A customer requested a GRE tunnel through the provider network between two customer sites using loopback to hide internal networks. Which configuration on R2 establishes the tunnel with R1?

- A. R2(config)# interface Tunnel 1R2(config-if)# ip address 172.20.1.2 255.255.255.0R2(config-if)# ip mtu 1400R2(config-if)# ip tcp adjust-mss 1360R2(config-if)# tunnel source 192.168.20.1R2(config-if)# tunnel destination 192.168.10.1
- B. R2(config)# interface Tunnel 1R2(config-if)# ip address 172.20.1.2 255.255.255.0R2(config-if)# ipmtu 1400R2(config-if)# ip tcp adjust-mss 1360R2(config-if)# tunnel source 10.10.2.2R2(config-if)# tunnel destination 10.10.1.1
- C. R2(config)# interface Tunnel 1R2(config-if)# ip address 172.20.1.2 255.255.255.0R2(config-if)# ipmtu 1500R2(config-if)# ip tcp adjust-mss 1360R2(config-if)# tunnel source 192.168.20.1R2(config-if)# tunnel destination 10.10.1.1
- D. R2(config)# interface Tunnel 1R2(config-if)# ip address 172.20.1.2 255.255.255.0R2(config-if)# ipmtu 1500R2(config-if)# ip tcp adjust-mss 1360R2(config-if)# tunnel source 10.10.2.2R2(config-if)# tunnel destination 10.10.1.1

CertEmpire

Answer:

B

Explanation:

To establish a GRE tunnel, the tunnel source and tunnel destination commands must use IP addresses that are routable over the transport network (the provider network). The question specifies using loopbacks, so R2's configuration must use its loopback (10.10.2.2) as the source and R1's loopback (10.10.1.1) as the destination.

Furthermore, GRE encapsulation adds a 24-byte header to the original IP packet. To prevent fragmentation over a standard 1500-byte MTU path, the tunnel interface's IP MTU must be reduced. A value of 1400 is a common and safe practice. The ip tcp adjust-mss 1360 command complements this by setting the TCP Maximum Segment Size to 1360 (1400 MTU - 40 bytes for IP/TCP headers), further preventing fragmentation for TCP traffic.

Why Incorrect Options are Wrong:

- A. The tunnel source and tunnel destination use private IP addresses (192.168.x.x) which are internal to the customer sites and not routable across the provider network.
- C. This option incorrectly uses a private, non-routable IP address (192.168.20.1) for the tunnel source and an incorrect ip mtu 1500, which will cause fragmentation.
- D. While this option correctly identifies the source and destination loopback addresses, it incorrectly sets the ip mtu to 1500. This will cause fragmentation as it does not account for the 24-byte GRE overhead.

References:

1. Cisco Systems, "Generic Routing Encapsulation (GRE) Configuration Guide, Cisco IOS XE Gibraltar 16.12.x". In the "Configuring a GRE Tunnel" section, the documentation specifies that the tunnel source and tunnel destination must be configured with addresses reachable over the transport network. It also notes that the IP MTU on the tunnel is automatically set to 24 bytes less than the outgoing interface MTU to account for the GRE header, making a manual setting of 1500 incorrect for a standard Ethernet path.
2. Cisco Systems, "Resolve IP Fragmentation, MTU, MSS, and PMTUD Issues with GRE and IPsec". This document explicitly states, "Because of the overhead of the GRE tunnel, you might need to reduce the MTU on the tunnel interface. A setting of 1400 bytes is a good starting point for the MTU on the tunnel interface." It also explains that `ip tcp adjust-mss` should be set to the tunnel MTU minus 40 bytes (e.g., $1400 - 40 = 1360$), which directly validates the values used in the correct answer.
3. RFC 2784, "Generic Routing Encapsulation (GRE)". Section 2, "Packet Format," defines the GRE header, which, combined with the new IP header, typically adds 24 bytes of overhead, necessitating a reduction in the effective MTU for the encapsulated packet.

Question: 2

A network administrator added a new spoke site with dynamic IP on the DMVPN network. Which configuration command passes traffic on the DMVPN tunnel from the spoke router?

- A. ip nhrp registration ignore
- B. ip nhrp registration no-registration
- C. ip nhrp registration dynamic
- D. ip nhrp registration no-unique

Answer:

D

Explanation:

This command is configured on the DMVPN hub, not the spoke router as the question incorrectly implies. It is required when multiple spokes are behind the same NAT/PAT device and thus share a single public (NBMA) IP address. By default, the hub's NHRP database requires a unique NBMA-to-tunnel IP mapping. The no-unique keyword disables this check, allowing the new dynamic spoke to successfully register with the hub even if its public IP is already in use by another spoke. This successful registration is a prerequisite for traffic to pass through the tunnel.

CertEmpire

Why Incorrect Options are Wrong:

A. ip nhrp registration ignore

This hub command ignores all dynamic NHRP registrations, which is the opposite of what is needed to add a new dynamic spoke.

B. ip nhrp registration no-registration

This spoke command disables the sending of registration packets, preventing the hub from ever learning the spoke's dynamic IP address.

C. ip nhrp registration dynamic

This is not a valid Cisco IOS command for NHRP configuration; it is a syntactically incorrect option.

References:

1. Cisco Systems, IP Addressing: NHRP Command Reference.

Reference for D: Under the ip nhrp registration no-unique command description, it states, "To allow the hub to overwrite existing Next Hop Resolution Protocol (NHRP) entries, use the ip nhrp registration no-unique command... This command is useful when a spoke has a dynamic address and is behind a Port Address Translation (PAT) device." This confirms the command's purpose and its application on the hub for dynamic spokes.

<https://certempire.com>

Reference for A & B: The same document details `ip nhrp registration ignore` as a hub command to reject dynamic registrations and `ip nhrp registration no-registration` as a spoke command to disable registrations.

2. Cisco Systems, DMVPN Configuration Guide, Cisco IOS XE Release 3S, "DMVPN with NAT" section.

Reference for D: This guide explains the challenge of spokes behind a PAT device: "The hub needs to differentiate the spokes although they have the same IP address." It then provides a configuration example for the hub that includes the `ip nhrp registration no-unique` command to solve this exact issue.

3. Cisco Systems, Implementing Secure Converged Wide Area Networks (ISCW) Version 3.0, Student Guide, Volume 2, Module 7, "Implementing DMVPN".

Reference for D: In the section "DMVPN and NAT," this official courseware explains, "The hub router must be configured with the `ip nhrp registration no-unique` command. This command allows the hub to accept multiple NHRP registrations for the same NBMA address." This directly supports the use of option D in a scenario with dynamic spokes that may be subject to NAT.

Question: 3

Which IPv6 feature enables a device to reject traffic when it is originated from an address that is not stored in the device binding table?

- A. IPv6 Snooping
- B. IPv6 Source Guard
- C. IPv6 DAD Proxy
- D. IPv6 RA Guard

Answer:

B

Explanation:

IPv6 Source Guard is a security feature that filters traffic based on the source IPv6 address and, optionally, the source MAC address. It leverages the binding table, which is populated by the IPv6 Snooping feature, to validate incoming packets. When IPv6 Source Guard is enabled on an interface, the switch examines traffic from hosts connected to that interface. If the source address of an incoming packet does not match a valid entry in the binding table for that specific port, the packet is dropped. This effectively prevents IP address spoofing attacks from devices on that segment.

CertEmpire

Why Incorrect Options are Wrong:

- A. IPv6 Snooping: This feature builds and maintains the binding table by observing Neighbor Discovery Protocol messages, but it does not, by itself, perform the traffic filtering or rejection.
- C. IPv6 DAD Proxy: The Duplicate Address Detection (DAD) Proxy feature responds to DAD queries on behalf of hosts, which helps in scaling ND-P in large L2 domains. It is not a traffic filtering mechanism.
- D. IPv6 RA Guard: This feature specifically filters and drops unauthorized Router Advertisement (RA) messages from rogue devices. It does not filter general data traffic based on the source address.

References:

1. Cisco Systems, Inc. (2023). First-Hop Security Configuration Guide, Cisco IOS XE Bengaluru 17.6.x (Catalyst 9300 Switches). Chapter: "Configuring IPv6 First-Hop Security", Section: "IPv6 Source Guard".

The document states, "IPv6 source guard is a security feature that filters traffic based on the IPv6 source address and the source MAC address... It uses the binding table populated by IPv6 snooping to validate the source of the incoming packets. When IPv6 source guard is enabled on an interface, the device drops packets when the source address is not found in the binding table."

<https://certempire.com>

2. Cisco Systems, Inc. (2021). IP Addressing: IPv6 Configuration Guide, Cisco IOS XE 17. Chapter: "Implementing IPv6 First-Hop Security", Section: "IPv6 Source Guard".

This guide explains, "The IPv6 Source Guard feature filters traffic based on the IPv6 source address of the traffic. The feature uses the binding table, which is populated by the IPv6 snooping feature, to validate the source of the incoming packets."

3. Haddad, W., & Bshara, M. (2018). IPv6 Security. Cisco Press. Chapter 4: "Layer 2 Security".

This official Cisco Press publication details that IPv6 Source Guard is the feature responsible for using the binding table to perform source IP address filtering, thereby preventing spoofing attacks. It explicitly differentiates its role from IPv6 Snooping, which builds the table, and RA Guard, which protects against rogue routers.

CertEmpire

Question: 4

Refer to the exhibit.

```
R1#sh ip route
Codes: L - local, C - connected, S - static, R - RIP, M - mobile, B - BGP
       D - EIGRP, EX - EIGRP external, O - OSPF, IA - OSPF inter area
       N1 - OSPF NSSA external type 1, N2 - OSPF NSSA external type 2
       E1 - OSPF external type 1, E2 - OSPF external type 2
       i - IS-IS, su - IS-IS summary, L1 - IS-IS level-1, L2 - IS-IS level-2
       ia - IS-IS inter area, * - candidate default, U - per-user static route
       o - ODR, P - periodic downloaded static route, H - NHRP, l - LISP
       a - application route
       + - replicated route, % - next hop override, p - overrides from PfR

Gateway of last resort is not set

D    10.0.0.0/8 [90/409600] via 172.16.1.200, 00:00:28, Ethernet0/0
    172.16.0.0/16 is variably subnetted, 2 subnets, 2 masks
C    172.16.1.0/24 is directly connected, Ethernet0/0
L    172.16.1.100/32 is directly connected, Ethernet0/0
    192.168.1.0/24 is variably subnetted, 2 subnets, 2 masks
C    192.168.1.0/24 is directly connected, Loopback0
L    192.168.1.100/32 is directly connected, Loopback0
R1#
```

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

Answer:

C

Explanation:

By default, routers use Administrative Distance (AD) to select the best path when multiple routing protocols provide a route to the same destination. The route with the lowest AD is preferred. The default AD for internal EIGRP is 90, and for RIP is 120. Therefore, R1 will initially prefer the EIGRP route.

To force R1 to use the RIP route, the AD of the EIGRP route must be made higher than the AD of the RIP route. The command `distance eigrp 130 120` under the EIGRP process on R1 changes the AD for internal EIGRP routes to 130. With this change, the RIP route's AD of 120 becomes more preferable than the EIGRP route's new AD of 130, causing R1 to install the RIP route.

Why Incorrect Options are Wrong:

- A. Setting RIP's AD to 90 creates a tie with EIGRP's default AD (90), which does not guarantee the RIP route will be chosen.
- B. Setting RIP's AD to 100 is still higher than EIGRP's default AD of 90, so the EIGRP route would remain preferred.
- D. Setting EIGRP's AD to 120 creates a tie with RIP's default AD (120), which does not guarantee the RIP route will be chosen.

References:

1. Cisco Systems, "IP Routing: Protocol-Independent Command Reference - distance (ip)": This document details the distance command. It confirms the default AD for RIP is 120. It also shows the syntax for modifying the AD of a specific routing protocol.

Reference Location: Under the command description for distance.

2. Cisco Systems, "IP Routing: EIGRP Configuration Guide - EIGRP Administrative Distance": This guide explains how to modify EIGRP's AD using the distance eigrp command. It specifies that the first value applies to internal routes and the second to external routes.

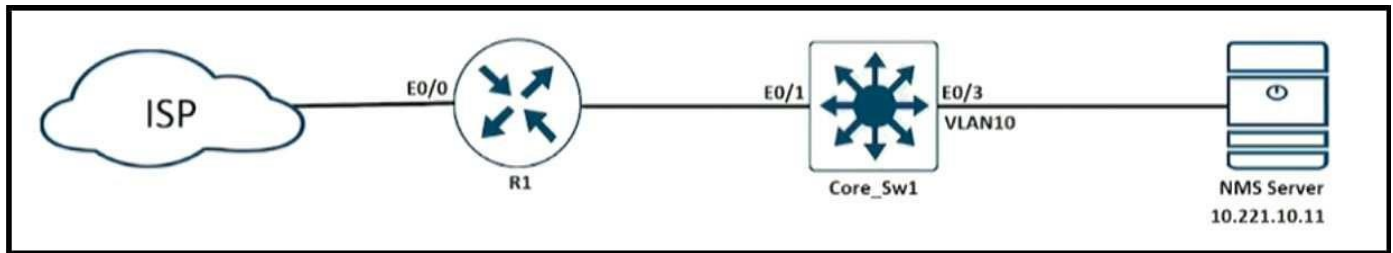
Reference Location: Section on "Configuring EIGRP," subsection "EIGRP Administrative Distance."

3. Cisco Systems, "What Is Administrative Distance?": This foundational document provides a table of default AD values. It explicitly lists "Enhanced Interior Gateway Routing Protocol (EIGRP)" with an AD of 90 and "Routing Information Protocol (RIP)" with an AD of 120.

Reference Location: The table of "Default Administrative Distance Values."

Question: 5

Refer to the exhibit.



During ISP router maintenance, the network produced many alerts because of the flapping interface. Which configuration on R1 resolves the issue?

- A. no snmp trap link-status
- B. snmp trap link-status down
- C. snmp trap ip verify drop-rate
- D. ip verify drop-rate notify hold-down 60

Answer:

B

CertEmpire

Explanation:

Interface flapping generates both ifDown and ifUp SNMP traps.

Applying the interface-level command

`snmp trap link-status down`

tells IOS to send only the link-down (ifDown) notification and suppress all link-up (ifUp) messages.

Therefore the NMS receives a single trap for every flap instead of two, eliminating the alert storm observed during the ISP's maintenance while still preserving the critical alarm that the link is down.

Why Incorrect Options are Wrong:

- A. no snmp trap link-status - Disables all link-status traps; removes the essential down alarm, not just the extra noise.
- C. snmp trap ip verify drop-rate - Controls uRPF drop-rate notifications, unrelated to physical link flaps.
- D. ip verify drop-rate notify hold-down 60 - Adds a hold-down timer for uRPF drop-rate traps, has no effect on interface link-status traps.

References:

1. Cisco IOS Interface and Hardware Component Command Reference, Release 15.5M&T - "snmp trap link-status" command, Description & Usage Notes, p. 1860: allows up down updown; default is updown.
2. Cisco IOS XE 17.3 SNMP Configuration Guide - section "Limiting LinkUp/LinkDown Notifications," Example 3: using 'snmp trap link-status down' to reduce trap volume.
3. Cisco IOS Security Command Reference, Release 15.5M&T - "ip verify drop-rate notify" command, p. 2478: applies to uRPF drop-rate events, not interface status.

CertEmpire

Question: 6

Refer to the exhibit.

```
ip vrf CCNP
  rd 1:1
interface Ethernet1
  ip vrf forwarding CCNP
  ip address 10.1.1.1 255.255.255.252
!
interface Ethernet2
  ip vrf forwarding CCNP
  ip address 10.2.2.2 255.255.255.252
```

Which configuration enables OSPF for area 0 interfaces to adjacency with a neighboring router with the same VRF?

- A. router ospf 1 vrf CCNPinterface Ethernet1ip ospf 1 area 0.0.0.0interface Ethernet2ip ospf 1 area 0.0.0.0
- B. router ospf 1interface Ethernet1ip ospf 1 area 0.0.0.0interface Ethernet2ip ospf 1 area 0.0.0.0
- C. router ospf 1 vrf CCNPnetwork 10.1.1.1 0.0.0.0 area 0network 10.2.2.2 0.0.0.0 area 0
- D. router ospf 1 vrf CCNPnetwork 10.0.0.0 0.0.255.255 area 0

Answer:

C

Explanation:

To enable OSPF within a specific Virtual Routing and Forwarding (VRF) instance, the OSPF process must be explicitly associated with that VRF. The command `router ospf 1 vrf CCNP` correctly initiates an OSPF process and links it to the "CCNP" VRF. Subsequently, the `network` command is used to enable OSPF on interfaces. Option C uses host-specific network statements

(10.1.1.1 0.0.0.0 and 10.2.2.2 0.0.0.0) which precisely match the IP addresses of Ethernet1 and Ethernet2. This is a correct and explicit method to enable OSPF only on the desired interfaces within the specified VRF and assign them to area 0.

Why Incorrect Options are Wrong:

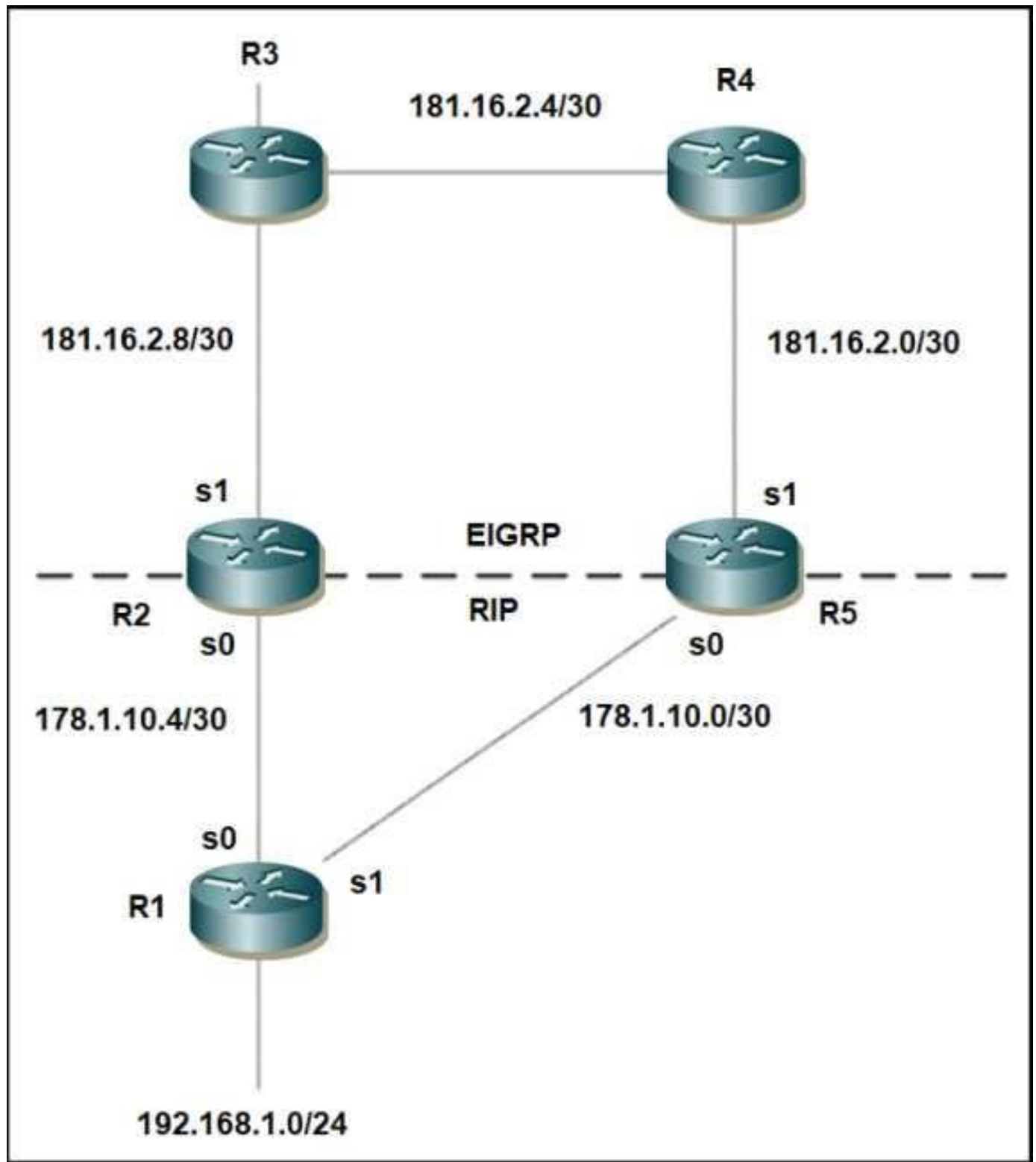
- A. While this is also a technically valid method using interface-level commands, option C represents the correct application of the router-level network command, which is equally valid and specific.
- B. This configuration creates an OSPF process in the global routing table, not within the "CCNP" VRF. The interfaces, being in the VRF, will not participate in this global OSPF process.
- D. This configuration uses a broad wildcard mask (0.0.255.255), which is less precise. It could unintentionally enable OSPF on other interfaces within the 10.0.0.0/16 range if they exist in the same VRF.

References:

1. Cisco Systems, Inc. (2022). IP Routing: OSPF Configuration Guide, Cisco IOS XE Amsterdam 17.3.x. Section: "Configuring OSPF", Subsection: "OSPFv2 VRF-Lite". This guide details the router ospf process-id vrf vrf-name command to create a VRF-aware OSPF process. It also provides examples using the network ip-address wildcard-mask area area-id command to enable OSPF on interfaces within that VRF.
2. Cisco Systems, Inc. (2022). Cisco IOS IP Routing: OSPF Command Reference. "network (OSPF)" command documentation. This reference explains that the network command under the router configuration mode defines on which interfaces OSPF will run. Using a 0.0.0.0 wildcard mask matches a specific host address.
3. Cisco Systems, Inc. (2022). Cisco IOS IP Routing: OSPF Command Reference. "router ospf" command documentation. This reference specifies the syntax router ospf process-id vrf vrf-name to enter router configuration mode for a specific OSPF process, optionally associating it with a VRF instance.

Question: 7

Refer to the exhibit.



Mutual redistribution is enabled between RIP and EIGRP on R2 and R5. Which configuration resolves the routing loop for the 192.168.1.0/24 network?

A. R2:router eigrp 10network 181.16.0.0redistribute rip metric 1 1 1 1 1distribute-list 1 in s1!router

```

ripnetwork 178.1.0.0redistribute eigrp 10 metric 2!access-list 1 deny 192.168.1.0access-list 1
permit

```

```

anyR5:router eigrp 10network 181.16.0.0redistribute rip metric 1 1 1 1 1distribute-list 1 in
s0!router

```

```

ripnetwork 178.1.0.0redistribute eigrp 10 metric 2!access-list 1 deny 192.168.1.0access-list 1
permit

```

```

any

```

B. R2:router eigrp 10network 181.16.0.0redistribute rip metric 1 1 1 1 1distribute-list 1 in s0!router
 ripnetwork 178.1.0.0redistribute eigrp 10 metric 2!access-list 1 deny 192.168.1.0access-list 1
 permit

```

anyR5:router eigrp 10network 181.16.0.0redistribute rip metric 1 1 1 1 1distribute-list 1 in
s0!router

```

```

ripnetwork 178.1.0.0redistribute eigrp 10 metric 2!access-list 1 deny 192.168.1.0access-list 1
permit

```

```

any

```

C. R2:router eigrp 10network 181.16.0.0redistribute rip metric 1 1 1 1 1distribute-list 1 in s0!router
 ripnetwork 178.1.0.0redistribute eigrp 10 metric 2!access-list 1 deny 192.168.1.0access-list 1
 permit

```

anyR5:router eigrp 10network 181.16.0.0redistribute rip metric 1 1 1 1 1distribute-list 1 in
s1!router

```

CertEmpire

```

ripnetwork 178.1.0.0redistribute eigrp 10 metric 2!access-list 1 deny 192.168.1.0access-list 1
permit

```

```

any

```

D. R2:router eigrp 7network 181.16.0.0redistribute rip metric 1 1 1 1 1distribute-list 1 in s1!router
 ripnetwork 178.1.0.0redistribute eigrp 7 metric 2!access-list 1 deny 192.168.1.0access-list 1
 permit

```

anyR5:router eigrp 7network 181.16.0.0redistribute rip metric 1 1 1 1 1distribute-list 1 in s1!router
ripnetwork 178.1.0.0redistribute eigrp 7 metric 2!access-list 1 deny 192.168.1.0access-list 1
permit

```

```

any

```

Answer:

A

Explanation:

The scenario describes a classic routing loop problem caused by mutual redistribution between two routing protocols (RIP and EIGRP) at two different points (R2 and R5). The 192.168.1.0/24 network originates in the RIP domain. It is redistributed into EIGRP by both R2 and R5. Consequently, R2 can learn this route from R5 via EIGRP, and vice-versa. This creates a

feedback loop.

To prevent this, we must filter the route from being learned back by the redistribution routers. The proposed solution uses an inbound distribute-list to prevent the EIGRP process on R2 and R5 from installing the 192.168.1.0/24 route from each other. According to the exhibit, R2's EIGRP interface is Serial1 (S1), and R5's EIGRP interface is Serial0 (S0). Therefore, the filter must be applied inbound on S1 for R2 and S0 for R5.

Why Incorrect Options are Wrong:

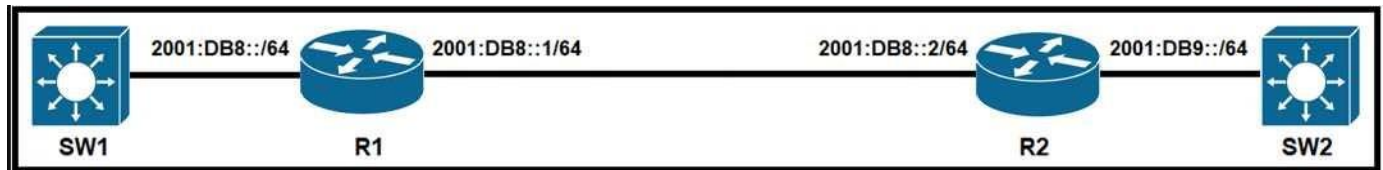
- B. The distribute-list on R2 is incorrectly applied to interface S0, which is in the RIP domain, not the EIGRP domain where the filter is needed.
- C. The distribute-lists on both R2 (S0) and R5 (S1) are applied to incorrect interfaces based on the network diagram.
- D. The distribute-list on R5 is incorrectly applied to interface S1. The exhibit clearly shows its EIGRP-facing interface is S0.

References:

1. Cisco IOS IP Routing: EIGRP Configuration Guide, "Filter Routes with a Distribute List". This document details the distribute-list in interface-type interface-number command syntax. It confirms that applying an inbound distribute-list on a specific interface filters routing updates received on that interface, which is the technique used in the correct answer.
2. Cisco IOS IP Routing: Protocol-Independent Configuration Guide, "Redistributing Routing Protocols". This guide discusses the challenges of redistribution, including routing loops. It states, "When you are redistributing routes, you can use distribute lists to prevent routing loops." This supports the fundamental strategy of using a distribute-list for loop prevention in this scenario.
3. Doyle, J., & Carroll, J. (2006). Routing TCP/IP, Volume 1 (2nd ed.). Cisco Press. Chapter 8, "Route Redistribution," discusses loop-prevention mechanisms. It explains that filtering redistributed routes is a primary method to avoid routing loops. Applying filters (like distribute-lists) on the routers performing redistribution prevents them from learning routes back from the other protocol domain. The configuration in option A correctly implements this principle.

Question: 8

Refer to the exhibit.



An engineer must advertise routes into IPv6 MP-BGP and failed. Which configuration resolves the issue on R1?

- A. `router bgp 65000no bgp default ipv4-unicastaddress-family ipv6 multicastnetwork 2001:DB8::/64`
- B. `router bgp 65000no bgp default ipv4-unicastaddress-family ipv6 unicastnetwork 2001:DB8::/64`
- C. `router bgp 64900no bgp default ipv4-unicastaddress-family ipv6 unicastnetwork 2001:DB8::/64`
- D. `router bgp 64900no bgp default ipv4-unicastaddress-family ipv6 multicastneighbor 2001:DB8:7000::2 translate-update ipv6 multicast`

Answer:

B

CertEmpire

Explanation:

The provided configuration on R1 (AS 65000) is missing the necessary commands to activate the IPv6 BGP session and advertise routes. To exchange IPv6 unicast prefixes, the neighbor must be activated within the address-family ipv6 unicast configuration mode. The network command is then used within this same address family to inject the specified prefix into the BGP table for advertisement. Option B correctly identifies the router's AS (65000) and places the network 2001:DB8::/64 command within the address-family ipv6 unicast stanza, which is the required procedure to advertise an IPv6 unicast route.

Why Incorrect Options are Wrong:

- A: This configures the network command under the ipv6 multicast address family, which is used for multicast routing, not standard unicast route advertisement.
- C: This configuration uses the incorrect autonomous system number (64900). The configuration must be applied on R1, which is in AS 65000.
- D: This uses the wrong AS number (64900) and specifies the incorrect address family (multicast) for advertising unicast reachability.

References:

1. Cisco IOS XE IP Routing: BGP Configuration Guide, "Implementing BGP for IPv6" section. This document outlines the required steps for IPv6 BGP configuration. It states, "The network command configured within the IPv6 address family configuration mode is used to advertise a prefix into the IPv6 BGP database." This confirms that the network command must be under address-family ipv6 unicast.

Source: Cisco Systems, "IP Routing: BGP Configuration Guide, Cisco IOS XE Fuji 16.9.x", Chapter: BGP for IPv6.

2. Cisco IOS IP Routing: BGP Command Reference, network (BGP) command. The documentation for the network command specifies its usage within an address family. For IPv6, it must be configured under address-family ipv6. The example provided shows: router bgp 65000 ... address-family ipv6 unicast ... network 2001:DB8:1::/48.

Source: Cisco Systems, "Cisco IOS IP Routing: BGP Command Reference", network (BGP) command documentation.

3. Cisco IOS IP Routing: BGP Configuration Guide, "BGP Support for the IPv6 Address Family" section. This guide explicitly shows that to advertise an IPv6 network, you must enter the IPv6 address family configuration mode and use the network command. It also confirms that neighbors must be activated within this address family using the neighbor activate command.

Source: Cisco Systems, "IP Routing: BGP Configuration Guide, Cisco IOS Release 15M&T", Section: IPv6 Routing: BGP.

CertEmpire

Question: 9

An engineer failed to run diagnostic commands on devices using Cisco DNA Center. Which action in Cisco DNA Center resolves the issue?

- A. Enable Command Runner
- B. Enable APIs
- C. Enable CDP
- D. Enable Secure Shell

Answer:

A

Explanation:

The Command Runner is the specific tool within the Cisco DNA Center GUI designed to execute CLI commands, including diagnostic commands, on managed network devices. If an engineer is unable to run these commands, the most direct cause related to a Cisco DNA Center feature is the inability to access or use this tool. The term "Enable Command Runner" in this context refers to providing the engineer with the necessary Role-Based Access Control (RBAC) permissions to use the tool, which would resolve the issue.

CertEmpire

Why Incorrect Options are Wrong:

- B. Enable APIs: APIs are for programmatic access and automation. The scenario describes an engineer interacting with the platform to run commands, which is a function of the GUI-based Command Runner tool.
- C. Enable CDP: Cisco Discovery Protocol (CDP) is used for device discovery and topology mapping. It is not involved in the execution of CLI commands on devices that are already managed.
- D. Enable Secure Shell: While SSH is the underlying protocol used by Cisco DNA Center to communicate with devices, the inability to run commands points to the specific tool, not the protocol itself.

References:

1. Cisco DNA Center User Guide, 2.3.5 - Command Runner: "You can use the Command Runner tool to run CLI commands on a single device or on multiple devices of the same type... You can run all commands, including configuration commands, on a device." This document confirms that Command Runner is the designated tool for this task.
2. Cisco DNA Center Administrator Guide, 2.3.5 - About Role-Based Access Control: This guide details how user roles (like NETWORK-ADMIN-ROLE and OBSERVER-ROLE) are granted specific privileges to access tools. For example, a user with an OBSERVER role may not have

<https://certempire.com>

permission to run commands. "Enabling" the feature for the engineer would involve assigning them a role with the appropriate permissions for the Command Runner tool.

CertEmpire

Question: 10

Which two components are required for MPLS Layer 3 VPN configuration? (Choose two)

- A. Use pseudowire for Layer 2 routes
- B. Use MP-BGP for customer routes
- C. Use OSPF between PE and CE
- D. Use a unique RD per customer VRF
- E. Use LDP for customer routes

Answer:

B, D

Explanation:

An MPLS Layer 3 VPN architecture fundamentally relies on two core components to function. First, a unique Route Distinguisher (RD) is required for each customer's Virtual Routing and Forwarding (VRF) instance. The RD is prepended to the customer's non-unique IPv4 prefix to create a globally unique 96-bit VPNv4 prefix. This process is essential for allowing different customers to use overlapping IP address spaces. Second, Multiprotocol BGP (MP-BGP) is the control plane protocol required to advertise these VPNv4 prefixes between Provider Edge (PE) routers across the service provider's core network. MP-BGP carries the customer routes along with their associated RDs and Route Targets (RTs).

Why Incorrect Options are Wrong:

- A. Pseudowires are the data plane mechanism used for Layer 2 VPNs (VPLS/EoMPLS), not Layer 3 VPNs.
- C. OSPF is only one of several optional choices for the PE-CE routing protocol; BGP, EIGRP, RIP, or static routes are also valid.
- E. LDP is used to exchange labels for routes within the provider's core (to build LSPs), not to distribute customer routes.

References:

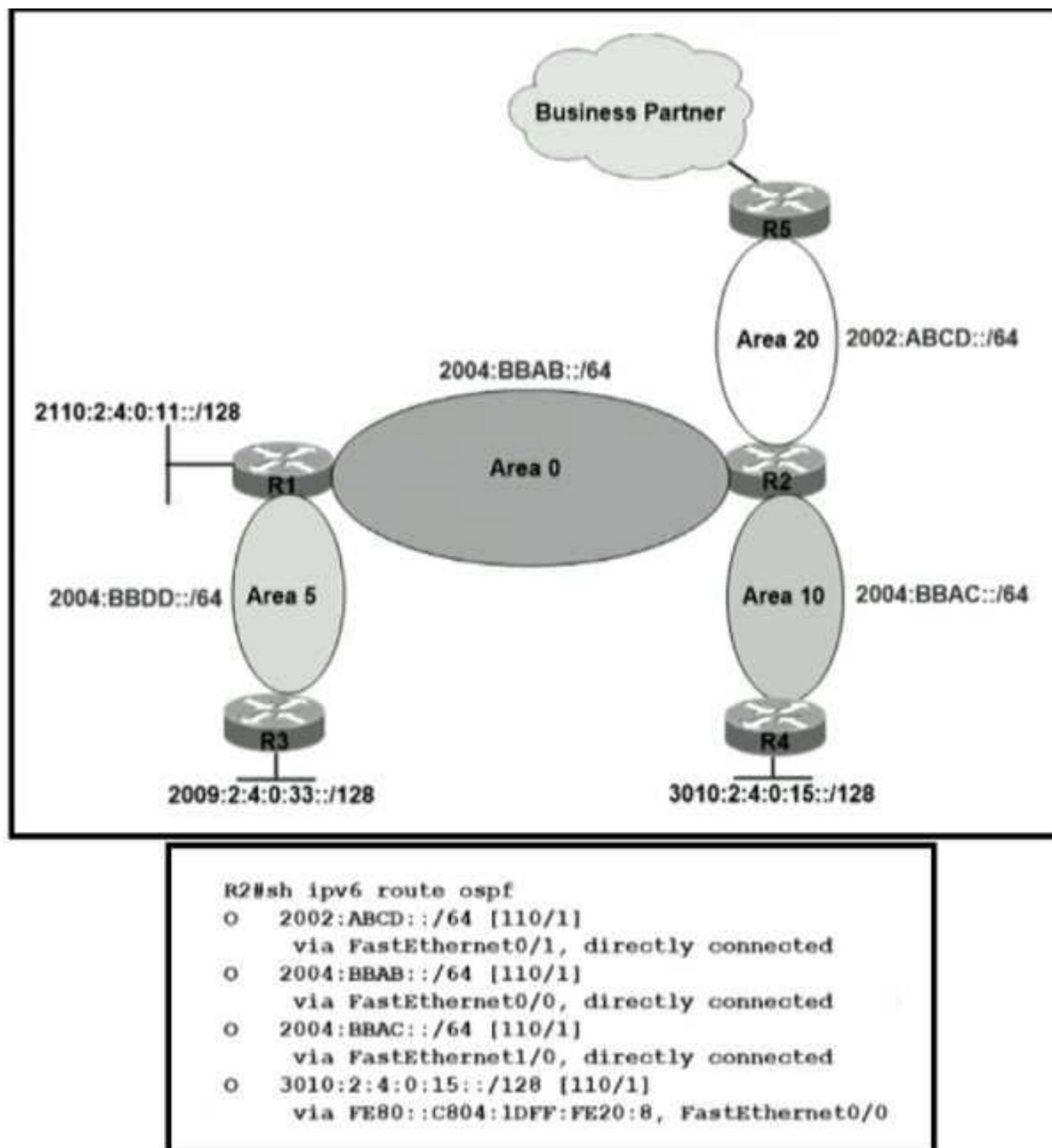
1. Cisco Systems, Inc., MPLS Layer 3 VPNs Configuration Guide, Cisco IOS XE Amsterdam 17.3.x, "MPLS Layer 3 VPNs Overview". This guide states, "MPLS VPNs use the BGP protocol to distribute VPN routing information across the provider's backbone." It also specifies, "A VPN route is a 12-byte quantity, beginning with an 8-byte route distinguisher (RD) and ending with a 4-byte IPv4 address prefix." This directly supports the requirement for MP-BGP and a unique RD.
2. Rosen, E., & Rekhter, Y., RFC 4364: BGP/MPLS IP Virtual Private Networks (VPNs), Internet Engineering Task Force (IETF), February 2006. Section 4, "Distributing VPN-IPv4 Routes

Between PEs," explicitly details the use of Multiprotocol BGP (MP-BGP) extensions. Section 4.2, "The VPN-IPv4 Address Family," defines the structure of the VPN-IPv4 address, which consists of an 8-byte Route Distinguisher (RD) and a 4-byte IPv4 address.

3. Cisco Press, Implementing and Operating Cisco Enterprise Network Core Technologies (ENCOR) v1.0 Official Cert Guide, Chapter 10: "MPLS". The section "MPLS Layer 3 VPN Components" identifies the key building blocks as VRFs, Route Distinguishers (RDs), Route Targets (RTs), and Multiprotocol BGP (MP-BGP). It clarifies that PE-CE routing protocols like OSPF are a design choice, not a universal requirement.

Question: 11

Refer to the exhibit.



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:

D

Explanation:

The objective is to prevent routes from Area 5 from being advertised into Area 20, while allowing them in all other areas. The initial filtering attempt on the Area 5 ABR (R1) blocked the routes from entering the backbone (Area 0), thus preventing their propagation to any other area. The correct location to apply the filter is on the Area Border Router (ABR) that injects routes into the target area. In this topology, R2 is the ABR for Area 20. By applying a filter list in the out direction for Area 20 on router R2, the advertisement of Type-3 LSAs (inter-area routes) from the backbone into Area 20 is controlled. This configuration correctly blocks the Area 5 routes from entering the business partner network (Area 20) without affecting their availability in Area 0 or other areas.

Why Incorrect Options are Wrong:

CertEmpire

- A. Applying a filter in the in direction on R1 for Area 5 would filter routes being installed in the RIB within Area 5, not routes being advertised out of it.
- B. This is the likely cause of the original problem. An out filter on R1 for Area 5 prevents its routes from being advertised as Type-3 LSAs into the backbone (Area 0).
- C. An in filter on R2 for Area 20 would filter routes from LSAs already present in Area 20's LSDB from being installed in the RIB, not prevent their advertisement into the area.

References:

1. Cisco IOS XE IP Routing: OSPF Configuration Guide, Cisco IOS XE Fuji 16.9.x, "OSPFv3 Interarea Route Filtering".
Section: How to Configure OSPFv3 Filtering OSPFv3 Interarea Routes
Content: The documentation for the area filter-list command states that the out keyword is used to "Filter routes advertised from other areas into the specified area." This confirms that to block routes from entering Area 20, the filter must be applied with the out keyword for Area 20 on its ABR (R2).
2. Cisco Press, Implementing Cisco Enterprise Advanced Routing and Services (ENARSI) v1.0 Official Cert Guide, by Raymond Lacoste and Brad Edgeworth.
Chapter/Section: Chapter 6, "OSPF," section "OSPF Route Filtering," subsection "Filtering on an ABR with a Filter List."

Content: The text explains that the `area area-id filter-list prefix-list-name out` command, when configured on an ABR, filters Type-3 LSAs from being advertised into the specified area. This directly supports applying the filter on R2 for Area 20 in the out direction to prevent routes from Area 0 (and by extension, Area 5) from entering Area 20.

3. Cisco IOS IP Routing: OSPF Command Reference, "area filter-list".

Section: Usage Guidelines

Content: This reference specifies that the `area filter-list` command is used for ABR-type filtering between OSPF areas. It reiterates that the `out` parameter filters routes being advertised from other areas into the specified area, which is the precise requirement of the scenario.

CertEmpire

Question: 12

Refer to the exhibit.

```
ipv6 dhcp pool DHCPPOOL
address prefix 2001:0:1:4::/64 lifetime infinite infinite

interface FastEthernet0/0
ip address 10.0.0.1 255.255.255.240
duplex auto
speed auto
ipv6 address 2001:0:1:4::1/64
ipv6 enable
ipv6 nd ra suppress
ipv6 ospf 1 area 1
ipv6 dhcp server DHCPPOOL
```

Reachability between servers in a network deployed with DHCPv6 is unstable. Which command must be removed from the configuration to make DHCPv6 function?

- A. ipv6 dhcp server DHCPPOOL
- B. ipv6 address 2001:0:1:4::/64
- C. ipv6 nd ra suppress
- D. address prefix 2001:0:1:4::/64 lifetime infinite infinite

Answer:

C

Explanation:

For a client to use stateful DHCPv6 to obtain an IPv6 address, it must be notified to do so by the local router. This notification is delivered via Router Advertisement (RA) messages containing a "Managed Address Configuration" (M) flag set to 1. The command `ipv6 nd ra suppress` completely stops the router from sending any RA messages on the interface. Without receiving an RA message with the M flag set, clients on the network segment will not initiate the DHCPv6 process (SOLICIT, ADVERTISE, etc.) to acquire an address, thus breaking the DHCPv6 service and

causing unstable reachability. Removing this command is necessary to restore DHCPv6 functionality.

Why Incorrect Options are Wrong:

- A. `ipv6 dhcp server DHCPPOOL`: This command is essential; it enables the DHCPv6 server process on the interface. Removing it would disable the service.
- B. `ipv6 address 2001:0:1:4::/64`: This command configures the router's own address on the interface, which is required for it to operate on that IPv6 segment.
- D. `address prefix 2001:0:1:4::/64 lifetime infinite infinite`: This command is configured within the DHCPv6 pool, not the interface, and is necessary to define the address range for allocation.

References:

1. Cisco Systems, Inc. (2023). Cisco IOS XE Dublin 17.12.x (Catalyst 9300 Switches) IP Addressing: IPv6 Configuration Guide.
Section: "DHCP for IPv6" "DHCPv6 Address Assignment"
Content: The documentation states, "A host can acquire IPv6 addresses in the following ways: ... Stateful autoconfiguration using DHCPv6... The M flag indicates whether to use DHCPv6 to get a stateful address." This establishes the dependency on flags within RA messages for stateful DHCPv6.
2. Cisco Systems, Inc. (2023). Cisco IOS XE Dublin 17.12.x (Catalyst 9300 Switches) IPv6 Command Reference.
Section: "ipv6 nd" commands `ipv6 nd ra suppress`
Content: The command description explicitly states its function: "To suppress the sending of router advertisements on an interface, use the `ipv6 nd ra suppress` command in interface configuration mode." This directly confirms that the command in question prevents the necessary RA messages from being sent.
3. Hinden, R., & Deering, S. (2006). RFC 4861: Neighbor Discovery for IP version 6 (IPv6).
Section: 6.2.3. "Router Advertisement Message Content"
Content: This standard defines the "Managed address configuration" flag (M flag). It specifies that "If set, it indicates that addresses are available via Dynamic Host Configuration Protocol DHCPv6." This foundational document confirms that the M flag in an RA is the standard mechanism to trigger stateful DHCPv6 client behavior. Suppressing the RA prevents this trigger.

Question: 13

Refer to the exhibit.

```
ip prefix-list DMZ-STATIC seq 5 permit 10.1.1.0/24
!
route-map DMZ permit 10
    match ip address prefix-list DMZ-STATIC
!
router ospf 1
network 0.0.0.0 0.0.0.0 area 0
redistribute static route-map DMZ
!
ip route 10.1.1.0 255.255.255.0 10.20.20.1
```

The static route is not present in the routing table of an adjacent OSPF neighbor router. Which action resolves the issue?

- A. Configure the next hop of 10.20.20.1 in the prefix list DMZ-STATIC
- B. Configure the next-hop interface at the end of the static route for it to get redistributed
- C. Configure a permit 20 statement to the route map to redistribute the static route
- D. Configure the subnets keyword in the redistribution command

Answer:

D

Explanation:

By default, when redistributing static routes into OSPF, the router only considers classful networks (e.g., 10.0.0.0/8, 172.16.0.0/16). The static route configured, 10.10.10.0/24, is a subnet of the class A network 10.0.0.0/8. For OSPF to redistribute this and other subnetted static routes, the subnets keyword must be explicitly included in the redistribute command. The route-map and prefix-list are correctly configured to filter for this specific route, but the route is never considered for redistribution in the first place without the subnets keyword. The correct command is redistribute static subnets route-map DMZ-STATIC.

Why Incorrect Options are Wrong:

- A. A prefix-list is used to match network prefixes (address and mask), not the next-hop address of a route.
- B. A static route is eligible for redistribution whether it is configured with a next-hop IP address or

an exit interface.

C. The existing permit 10 statement in the route-map already correctly matches the prefix-list and permits the route.

References:

1. Cisco Systems, Inc., IP Routing: OSPF Configuration Guide, Cisco IOS XE Gibraltar 16.12.x, "How to Redistribute Routing Information" section. The guide explains that to redistribute routes from protocols that support variable-length subnet masking (VLSM), such as static routes for subnets, the subnets keyword is required.
2. Cisco Systems, Inc., Cisco IOS IP Routing: OSPF Command Reference, "redistribute (OSPF)" command documentation. The syntax is shown as redistribute ... subnets. The description for the subnets keyword states, "(Optional) Redistributes subnetted routes. Without this keyword, only routes that are not subnetted are redistributed." This directly addresses the issue in the question.

CertEmpire

Question: 14

Refer to the exhibit.

```
!-- ACL for CoPP Routing class-map
!
access-list 120 permit tcp any gt 1024 eq bgp log
access-list 120 permit tcp any bgp gt 1024 established
access-list 120 permit tcp any gt 1024 eq 639
access-list 120 permit tcp any eq 639 gt 1024 established
access-list 120 permit tcp any eq 646
access-list 120 permit udp any eq 646
access-list 120 permit ospf any
access-list 120 permit ospf any host 224.0.0.5
access-list 120 permit ospf any host 224.0.0.6
access-list 120 permit eigrp any
access-list 120 permit eigrp any host 224.0.0.10
access-list 120 permit udp any any eq pim-auto-rp
```

The control plane is heavily impacted after the CoPP configuration is applied to the router. Which command removal lessens the impact on the control plane?

- A. access-list 120 permit udp any any eq pim-auto-rp
- B. access-list 120 permit eigrp any host 224.0.0.10
- C. access-list 120 permit ospf any
- D. access-list 120 permit tcp any gt 1024 eq bgp log

Answer:

A

Explanation:

The configuration applies a single, highly restrictive policer of 8000 bps to a class that includes four different protocols: EIGRP, OSPF, BGP, and PIM Auto-RP. A heavy impact on the control plane indicates that this rate is too low for the aggregate traffic, causing legitimate protocol packets to be dropped.

Among the listed protocols, EIGRP, OSPF, and BGP are fundamental unicast routing protocols essential for network stability. PIM Auto-RP, however, is used for multicast rendezvous point

discovery. While important for multicast services, its traffic can be voluminous, and it is generally less critical for core network reachability than the unicast protocols. Removing the PIM Auto-RP permit statement from the access list would exclude its traffic from this restrictive policer, freeing up the limited bandwidth for the essential unicast routing protocols and thus mitigating the impact.

Why Incorrect Options are Wrong:

- B. Removing the EIGRP entry would cause EIGRP adjacencies to drop, severely worsening the impact on the control plane and network connectivity.
- C. Removing the OSPF entry would break OSPF adjacencies, which is a critical failure and would increase, not lessen, the control plane impact.
- D. Removing the BGP entry would cause BGP peering sessions to fail, disrupting inter-domain or large-scale internal routing, a critical network function.

References:

1. Cisco Systems, "Control Plane Policing Implementation Best Practices" White Paper.
Reference: Section "Traffic Classification," Table 1, "Control Plane Traffic Categories."
Details: This document categorizes control plane traffic. While routing protocols like BGP, OSPF, and EIGRP are listed as "Critically Important," they are often separated into their own classes for granular control. The document emphasizes that "the aggregate of all routing protocol traffic should not overwhelm the CP." CertEmpire Lumping multiple protocols, including the potentially chatty PIM Auto-RP, under a single low-rate policer is a common misconfiguration. The best practice is to isolate and police traffic types appropriately, and removing the less-critical or problematic protocol from the oversubscribed class is a valid troubleshooting step.
2. Cisco IOS XE Security Configuration Guide, "Configuring Control Plane Policing."
Reference: Chapter: "Control Plane Policing," Section: "Traffic Classification for CoPP."
Details: The official documentation states, "The key to a successful CoPP implementation is a clear understanding of the traffic that is required by the control plane to maintain network stability and connectivity." It advises separating different types of control-plane traffic to apply appropriate policies. In this scenario, the unicast routing protocols (BGP, OSPF, EIGRP) are more critical for basic network stability than the multicast PIM Auto-RP protocol. Therefore, removing PIM from the critical, under-provisioned class is the logical choice to restore stability.

Question: 15

Refer to the exhibit.

```
snmp-server community Public RO 90
snmp-server community Private RW 90
R1#show access-list 90
Standard IP access list 90
  permit 10.11.110.11
  permit 10.11.111.12
```

Nov 6 06:45:11: %SNMP-3-AUTHFAIL: Authentication failure for SNMP req from host 10.11.110.12

Nov 6 06:45:12: %SNMP-3-AUTHFAIL: Authentication failure for SNMP req from host 10.11.110.12

A network administrator notices these console messages from host 10.11.110.12 originating from interface E1/0. The administrator considers this an unauthorized attempt to access SNMP on R1. Which action prevents the attempts to reach R1 E1/0?

- A. Configure IOS control plane protection using ACL 90 on interface E1/0
- B. Configure IOS management plane protection using ACL 90 on interface E1/0
- C. Create an inbound ACL on interface E1/0 to deny SNMP from host 10.11.110.12
- D. Add a permit statement including the host 10.11.110.12 into ACL 90

Answer:

C

Explanation:

The console message %SNMP-3-AUTHFAIL confirms that unauthorized SNMP request packets from host 10.11.110.12 are successfully reaching and being processed by the router's control plane. The goal is to prevent these attempts entirely. The most direct and effective method is to filter this traffic at the ingress interface before it reaches the CPU for processing. Applying an inbound Access Control List (ACL) on interface E1/0 to explicitly deny SNMP traffic (UDP port 161) from the source IP 10.11.110.12 will cause the packets to be dropped at the interface level. This prevents the packets from ever reaching the SNMP process, thus stopping the unauthorized attempts and the resulting log messages.

Why Incorrect Options are Wrong:

- A. Control Plane Protection is applied globally to the control-plane virtual interface to rate-limit traffic, not applied directly to a physical interface to block a specific host.
- B. Management Plane Protection restricts management access to specific interfaces (e.g., a loopback), not by applying a standard traffic-filtering ACL on a transit interface like E1/0.
- D. Adding a permit statement for host 10.11.110.12 into ACL 90 would authorize the SNMP access, which is the opposite of the administrator's intent to block it.

References:

1. Cisco IOS IP Application Services Configuration Guide, Release 15M&T, "Configuring IP Access Lists" chapter, "Applying IP Access Lists to Interfaces" section.

This document details the procedure for filtering traffic on an interface. It states, "After you create an access list, you must apply it to an interface... When the software receives a packet on an interface, the software checks the source address of the packet against the access list for that interface." This directly supports applying an inbound ACL to filter unwanted traffic.

2. Cisco Systems, "Guide to Harden Cisco IOS Devices", Document ID: 13601, "Control Plane Policing" section.

This guide explains that Control Plane Policing (CoPP) is "applied to the aggregate control plane interface using the service-policy command." This confirms that CoPP is a global control-plane mechanism and is not applied to a specific physical interface like E1/0 to block a single host.

3. Cisco IOS Network Management Configuration Guide, Release 12.4, "Configuring SNMP Support" chapter, "snmp-server community" command reference.

The documentation for the snmp-server community command specifies that the associated access list "defines the source IP addresses that are permitted to use the community string." This confirms that adding the host to ACL 90 would permit, not deny, access, making option D incorrect.

Question: 16

Refer to the exhibit.

```
CPE# ping 10.0.2.4
Type escape sequence to abort.
Sending 5, 100-byte ICMP Echos to 10.0.2.4, timeout is 2
seconds:
!!!!
Success rate is 100 percent (5/5), round-trip min/avg/max =
1/1/1 ms
CPE# copy flash:/packages.conf tftp://10.0.2.4/
Address or name of remote host [10.0.2.4]?
Destination filename [packages.conf]?
%Error opening tftp://10.0.2.4/packages.conf (Undefined error)
```

The administrator is trying to overwrite an existing file on the TFTP server that was previously uploaded by another router. However, the attempt to update the file fails. Which action resolves this issue?

- A. Make the packages.conf file executable by all on the TFTP server
- B. Make the packages.conf file writable by all on the TFTP server
- C. Make sure to run the TFTP service on the TFTP server
- D. Make the TFTP folder writable by all on the TFTP server

Answer:

B

Explanation:

The exhibit displays the file permissions for packages.conf as -rw-r--r--. This permission set grants write access (w) only to the file's owner, which is root. The read permission (r) is granted to the owner, the group, and all other users. For security, TFTP server daemons often run as a non-privileged user (e.g., nobody, tftp). When the second router attempts to overwrite the file, the TFTP daemon, acting as a non-root user, is denied access because it falls into the "others" category, which lacks write permissions. To resolve this, the file must be made writable for the user account the TFTP service is running under. Making the file writable by all (chmod a+w packages.conf) is a direct way to ensure the TFTP daemon has the necessary permissions to overwrite the file.

Why Incorrect Options are Wrong:

- A. Executable permissions are for running a file as a program and are not required to write or overwrite a data file via TFTP.
- C. The fact that a file was previously uploaded successfully indicates that the TFTP service is already running on the server.
- D. Directory permissions control the ability to create or delete files within it, but permissions on the file itself control the ability to modify its contents.

References:

1. Linux tftpd(8) man page: The manual for the Trivial File Transfer Protocol daemon (tftpd) often specifies the security model. For many implementations, it states that for a file to be written (or overwritten), it must already exist and be publicly writable. The permissions -rw-r--r-- do not meet the "publicly writable" criteria for the TFTP daemon process if it is not running as the root user.
2. Arpaci-Dusseau, R. H., & Arpaci-Dusseau, A. C. (2018). Operating Systems: Three Easy Pieces. Arpaci-Dusseau Books. In Chapter 39, "Files and Directories," Section 39.3, "Permission Bits," the text explains that the write bit (w) for a file controls the ability to modify its contents. A process must have effective write permission on the file inode to perform a write operation.
3. Cisco IOS XE System Management Configuration Guide, Cisco IOS XE Gibraltar 16.12.x: In the chapter "Managing Configuration Files," the section "Working with Configuration Files" details the use of TFTP to copy files to and from a network device. The guide assumes a properly configured TFTP server, where file and directory permissions on the server allow the requested operations (read or write). The failure described in the question points to a server-side permission misconfiguration.

Question: 17

Refer to the exhibit.

```
R2#show ip route
```

```
Gateway of last resort is not set
```

```
10.0.0.0/8 is variably subnetted, 12 subnets, 3 masks
```

```
C 10.1.3.0/30 is directly connected, FastEthernet0/1
```

```
C 10.1.2.0/30 is directly connected, FastEthernet0/0
```

```
C 10.1.1.0/30 is directly connected, FastEthernet1/0
```

```
O E2 10.19.0.0/24 [110/20] via 10.1.3.2, 00:02:04, FastEthernet0/1
```

```
D 10.55.13.0/24 (90/4096001 via 10.1.2.2. 00:01:00. FastEthernet0/0
```

```
D 10.37.100.0/24 (90/4096001 via 10.1.2.2. 00:01:00. FastEthernet0/0
```

```
C 10.100.10.0/29 is directly connected, FastEthernet2/0.10
```

```
D 10.55.72.0/24 (90/409600) via 10.1.2.2. 00:01:01. FastEthernet0/0
```

```
C 10.100.20.0/29 is directly connected. FastEthernet2/0.20
```

```
O E2 10.144.1.0/24 /110/201 via 10.1.3.2. 00:12:51. FastEthernet0/1
```

```
D 10.55.144.0/24 (90/4096001 via 10.1.2.2. 00:01:01. FastEthernet0/0
```

```
O E2 10.123.187.0/24 [110/20] via 10.1.3.2. 00:12:51, FastEthernet0/1
```

```
R2#sh ip eigrp topology
```

```
IP-EIGRP Topology Table for AS(100)/ID(10.100.20.2)
```

```
Codes: P - Passive, A - Active, U - Update, Q - Query, R - Reply,
```

```
r- reply Status, s - sia Status
```

```
P 10.1.3.0/30, 1 successors, FD is 281600 via Connected, FastEthernet0/1
```

```
P 10.1.2.0/30, 1 successors, FD is 281600 via Connected, FastEthernet0/0
```

```
P 10.1.1.0/30, 1 successors, FD is 28160 via Connected, FastEthernet1/0
```

```
P 10.55.13.0/24, 1 successors, FD is 409600 via 10.1.2.2 (409600/128256). FastEthernet0/0
```

```
P 10.37.100.0/24, 1 successors, FD is 409600 via 10.1.2.2 (409600/128256). FastEthernet0/0
```

```
P 10.55.72.0/24, 1 successors, FD is 409600 via 10.1.2.2 (409600/128256), FastEthernet0/0
```

```
P 10.55.144.0/24, 1 successors, FD is 409600 via 10.1.2.2 (409600/128256), FastEthernet0/0
```

```
P 10.123.187.0/24, 0 successors, FD is Inaccessible via 10.1.2.2 (409600/128256), FastEthernet0/0
```

Router R2 should be learning the route for 10.123.187.0/24 via EIGRP. Which action resolves the issue without introducing more issues?

- A. Use distribute-list to modify the route as an internal EIGRP route
- B. Redistribute the route in EIGRP with metric, delay, and reliability
- C. Use distribute-list to filter the external router in OSPF
- D. Remove route redistribution in R2 for this route in OSPF

Answer:

C

Explanation:

The issue is a suboptimal routing problem caused by route feedback in a mutual redistribution scenario. Router R2 learns the 10.123.187.0/24 prefix from R1 via EIGRP (Administrative Distance 90). R2 then redistributes this route into OSPF, which R3 learns. R3, in turn, advertises

this route back to R2 as an OSPF external route (Administrative Distance 110).

Because R2 has the OSPF route in its routing table instead of the superior EIGRP route, it indicates a routing feedback problem. The standard solution is to prevent the redistribution router (R2) from re-learning the prefixes it originally advertised into the other protocol. Applying an inbound distribute-list within the OSPF process on R2 to filter this specific external prefix prevents it from being installed in the routing table, allowing the preferred EIGRP route to be used.

Why Incorrect Options are Wrong:

A. Use distribute-list to modify the route as an internal EIGRP route

A distribute-list is used for filtering routes, not for changing a route's type from external to internal.

B. Redistribute the route in EIGRP with metric, delay, and reliability

This action concerns routes being redistributed into EIGRP, but the problem is with a route being learned from OSPF that should be learned from EIGRP.

D. Remove route redistribution in R2 for this route in OSPF

This would stop R3 and the rest of the OSPF domain from learning the route, breaking connectivity and introducing a new, more significant issue.

References:

1. Cisco Systems, Inc., IP Routing: OSPF Configuration Guide, Cisco IOS Release 15M&T, "How to Filter OSPF LSAs". The section on distribute-list in command explains its use for filtering routes from the OSPF routing table. This mechanism is used to prevent routing loops or suboptimal routing in redistribution scenarios.

Reference: Document ID OL-29422-01, Chapter: "Filtering OSPF LSAs", Section: "OSPF Inbound Route Filtering".

2. Cisco Systems, Inc., Implementing Cisco IP Routing (ROUTE) Foundation Learning Guide, "Chapter 8: Implementing Route Redistribution". This official learning material details the challenges of mutual redistribution, including routing feedback. It presents filtering with distribute-lists as a primary solution to prevent a router from learning its own redistributed prefixes back from another routing domain.

Reference: ISBN: 978-1-58705-882-0, Chapter 8, Section: "Preventing Routing Loops in Hub-and-Spoke Topologies".

3. Lacoste, R., & Edgeworth, B., Implementing Cisco Enterprise Advanced Routing and Services (ENARSI) v1.0 Official Cert Guide. This book, directly aligned with the 300-410 exam, covers redistribution pitfalls extensively.

Reference: ISBN: 978-0-13-526208-8, Chapter 10: "Route Redistribution", Section: "Controlling Redistribution with Route Filtering". This section explicitly describes using inbound distribute-lists on the redistribution router to solve suboptimal routing caused by feedback.

Question: 18

Refer to the exhibit.

<pre> R2#show ip eigrp neighbors IP-EIGRP neighbors for process 100 H Address Interface Hold Uptime SRTT RTO Q Seq (sec) (ms) Cnt Num 1 192.168.10.1 Ser1/0 12 00:00:39 1 5000 2 0 *Jan 1 15:40:21.295: %DUAL-5-NBRCHANGE: IP-EIGRP(0) 100: Neighbor 192.168.10.1 (Serial1/0) is down: retry limit exceeded *Jan 1 15:40:51.567: %DUAL-5-NBRCHANGE: IP-EIGRP(0) 100: Neighbor 192.168.10.1 (Serial1/0) is up: new adjacency *Jan 1 15:42:11.107: %DUAL-5-NBRCHANGE: IP-EIGRP(0) 100: Neighbor 192.168.10.1 (Serial1/0) is down: retry limit exceeded *Jan 1 15:42:14.879: %DUAL-5-NBRCHANGE: IP-EIGRP(0) 100: Neighbor 192.168.10.1 (Serial1/0) is up: new adjacency </pre>		<pre> R1#show ip eigrp neighbors IP-EIGRP neighbors for process 100 </pre>	
<u>R1 Configuration:</u> <pre> key chain cisco key 2 key-string abc ! interface Loopback0 ip address 10.10.1.1 255.255.255.0 ! interface Serial1/0 ip address 192.168.10.1 255.255.255.0 ip authentication mode eigrp 100 md5 ip authentication key-chain eigrp 100 cisco serial restart-delay 0 ! router eigrp 100 network 10.10.1.0 0.0.0.255 network 192.168.10.0 no auto-summary </pre>		<u>R2 configuration:</u> <pre> key chain cisco key 1 key-string 123 key 2 key-string abc ! interface Loopback0 ip address 10.10.2.2 255.255.255.0 ! interface Serial1/0 ip address 192.168.10.2 255.255.255.0 ip authentication mode eigrp 100 md5 ip authentication key-chain eigrp 100 cisco no fair-queue ! ! router eigrp 100 network 10.10.2.0 0.0.0.255 network 192.168.10.0 no auto-summary </pre>	

R1 and R2 are configured for EIGRP peering using authentication and the neighbors failed to come up. Which action resolves the issue?

- A. Configure a matching key-id number on both routers
- B. Configure a matching lowest key-id on both routers
- C. Configure a matching key-chain name on both routers
- D. Configure a matching authentication type on both router

Answer:

A

Explanation:

For EIGRP MD5 authentication to establish a neighbor relationship, several parameters must be correctly configured. The exhibit shows that both routers use the same key-chain name (EIGRPKEY), the same authentication mode (MD5), and the same key-string (CISCO). However, R1 is configured with key 1, while R2 is configured with key 2. The key ID number is included in

the EIGRP packet header. The receiving router uses this ID to find the corresponding key in its local key chain to validate the message. Since R1 and R2 do not have a common key ID, authentication fails, and the peering cannot be established. Configuring a matching key ID on both routers will resolve the issue.

Why Incorrect Options are Wrong:

- B. The issue is the absence of any matching key ID, not which specific key ID (such as the lowest) is used.
- C. The key-chain name is locally significant to the router and is not exchanged between neighbors; therefore, it does not need to match.
- D. The authentication type (MD5) is already configured and matches on both routers, as shown by the `ip authentication mode eigrp 1 md5` command.

References:

1. Cisco Systems, Inc. (2023). IP Routing: EIGRP Configuration Guide, Cisco IOS XE Bengaluru 17.6.x.

Section: "How to Configure EIGRP MD5 Authentication"

Reference: In the configuration steps for the key, the guide states: "The key number must be the same on the local router and its neighbor." This directly confirms that the key ID is the mismatched parameter causing the failure. The document is available at: <https://www.cisco.com/c/en/us/td/docs/ios-xml/ios/iprouteeigrp/configuration/xr-17-6/ire-xe-17-6-book/ire-md5-authentication.html>

2. Vachon, R., Graziani, R., Edgeworth, A., & Hucaby, D. (2020). CCNP and CCIE Enterprise Core ENCOR 350-401 Official Cert Guide. Cisco Press.

Chapter 6: "EIGRP"

Section: "EIGRP Authentication"

Reference: The text explains that the EIGRP packet includes the key ID, and the receiving router uses that same key ID to look up its local key. This process requires that "both routers have a key with the same key ID and the same key string." The exhibit clearly shows a mismatch in the key ID (1 vs. 2).

3. Teare, D., & Paquet, C. (2015). Implementing Cisco IP Routing (ROUTE) Foundation Learning Guide. Cisco Press.

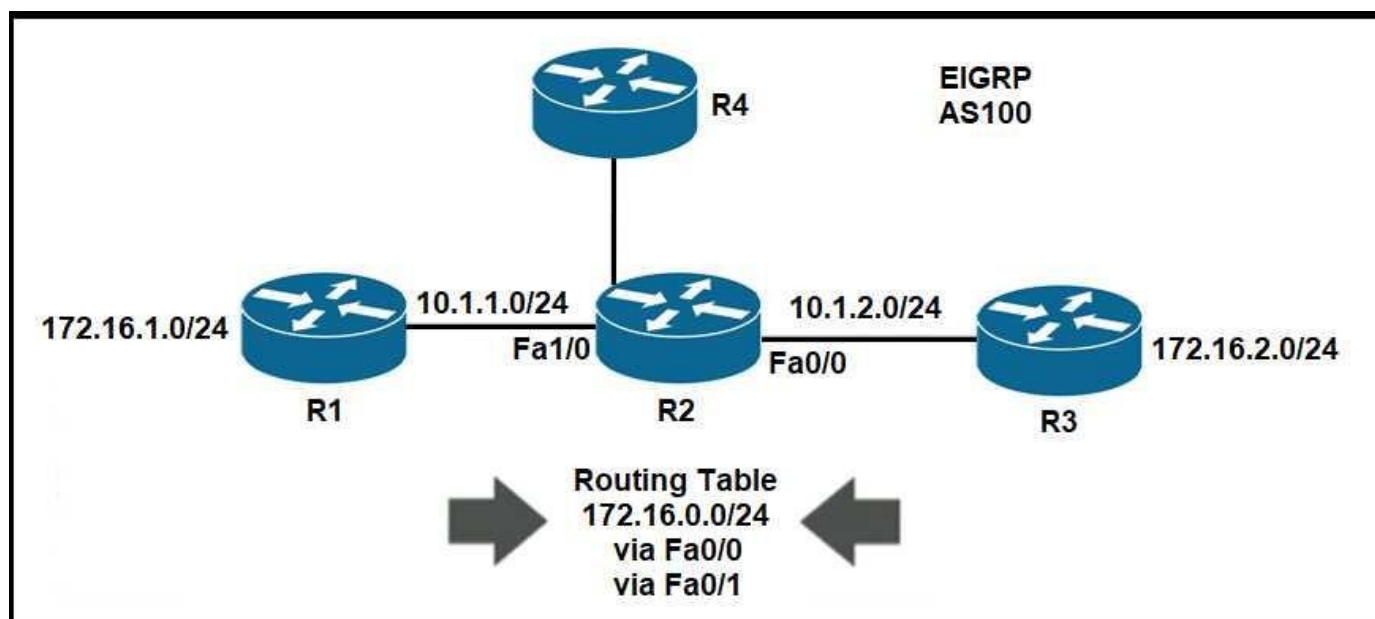
Chapter 3: "Implementing EIGRP"

Section: "EIGRP Authentication"

Reference: This guide clarifies that while the key-chain name is locally significant, the key number (ID) and the key string must match between neighbors for authentication to succeed. The configuration examples provided consistently show matching key IDs on peering routers.

Question: 19

Refer to the exhibit.



R4 is experiencing packet drop when trying to reach 172.16.2.7 behind R2. Which action resolves the issue?

- A. Insert a /16 floating static route on R2 toward R3 with metric 254
- B. Insert a /24 floating static route on R2 toward R3 with metric 254
- C. Enable auto summarization on all three routers R1, R2, and R3
- D. Disable auto summarization on R2

Answer:

D

Explanation:

The configuration output for R2 shows that auto-summary is enabled for EIGRP. This causes R2 to advertise its connected 172.16.0.0 subnets as a single classful summary route, 172.16.0.0/16, to its neighbors R1 and R3. R1 receives this summary and re-advertises it to R3.

As a result, R3's routing table shows two equal-cost paths to 172.16.0.0/16: one directly to R2 and another via R1. The path through R1 is invalid because R1's only route for this network is back to R2. When R3 load-balances traffic across these two paths, packets sent to R1 are dropped, creating a routing black hole. Disabling auto-summary on R2 with the no auto-summary command will cause it to advertise its more specific subnets, resolving the issue.

Why Incorrect Options are Wrong:

- A. A floating static route on R2 does not correct the invalid EIGRP route that R3 learns from R1.
- B. A floating static route on R2 does not correct the invalid EIGRP route that R3 learns from R1.
- C. Enabling auto-summary on all routers would likely introduce more summarization issues and would not fix the root cause.

References:

1. Cisco Systems, Inc., IP Routing: EIGRP Configuration Guide, Cisco IOS XE Gibraltar 16.12.x, "EIGRP Automatic Summary" section. This document states, "When automatic summarization is enabled, EIGRP summarizes on classful network boundaries." It further explains that disabling auto-summary is necessary in discontinuous networks to prevent routing issues.
2. Cisco Systems, Inc., IP Routing: EIGRP Command Reference, auto-summary (EIGRP) command documentation. The documentation specifies that the no auto-summary command is used "to disable this function and send subprefix routing information across classful network boundaries." This directly addresses the problem of advertising overly broad, classful summaries.
3. Graziani, R. (2015). IP Routing with EIGRP. Cisco Press. Chapter 6, "EIGRP Summarization," details how automatic summarization at classful boundaries can create routing black holes in discontinuous network topologies, similar to the one shown in the exhibit. The standard solution presented is to disable automatic summarization.

CertEmpire

Question: 20

Refer to the exhibit.

```
access-list 1 permit 209.165.200.215
access-list 2 permit 209.165.200.216
!
interface ethernet 1
ip policy route-map Texas
!
route-map Texas permit 10
match ip address 1
set ip precedence priority
set ip next-hop 209.165.200.217
!
route-map Texas permit 20
match ip address 2
set ip next-hop 209.165.200.218
```

Packets arriving from source 209.165.200.215 must be sent with the precedence bit set to 1, and packets arriving from source 209.165.200.216 must be sent with the precedence bit set to 5.

Which action resolves the issue?

- A. set ip precedence critical in route-map Texas permit 10
- B. set ip precedence critical in route-map Texas permit 20
- C. set ip precedence immediate in route-map Texas permit 10
- D. set ip precedence priority in route-map Texas permit 20

Answer:

B

Explanation:

The objective is to set IP Precedence to 1 for source 209.165.200.215 and 5 for source 209.165.200.216. The existing configuration in route-map Texas permit 10 correctly uses set ip precedence priority to set the value to 1 for the first source. The second source, 209.165.200.216, is matched by route-map Texas permit 20 but lacks a set command. According to Cisco documentation, the IP Precedence keyword critical corresponds to the numerical value 5. Therefore, adding set ip precedence critical to the permit 20 sequence completes the configuration as required.

Why Incorrect Options are Wrong:

- A. This incorrectly sets precedence to 5 (critical) for the source that requires precedence 1.
- C. This incorrectly sets precedence to 2 (immediate) for the source that requires precedence 1.
- D. This incorrectly sets precedence to 1 (priority) for the source that requires precedence 5.

References:

1. Cisco Systems, "IP Application Services Command Reference, Cisco IOS XE Gibraltar 16.12.x". This document details the set ip precedence (route-map) command and its arguments. It explicitly maps the keywords to their numerical values.

Reference: Under the "set ip precedence (route-map)" command description, the argument table lists critical as corresponding to precedence value 5, and priority to precedence value 1.

Link: <https://www.cisco.com/c/en/us/td/docs/ios-xml/ios/ipapp/command/iap-cr-book/iap-s1.html#wp3039189337>

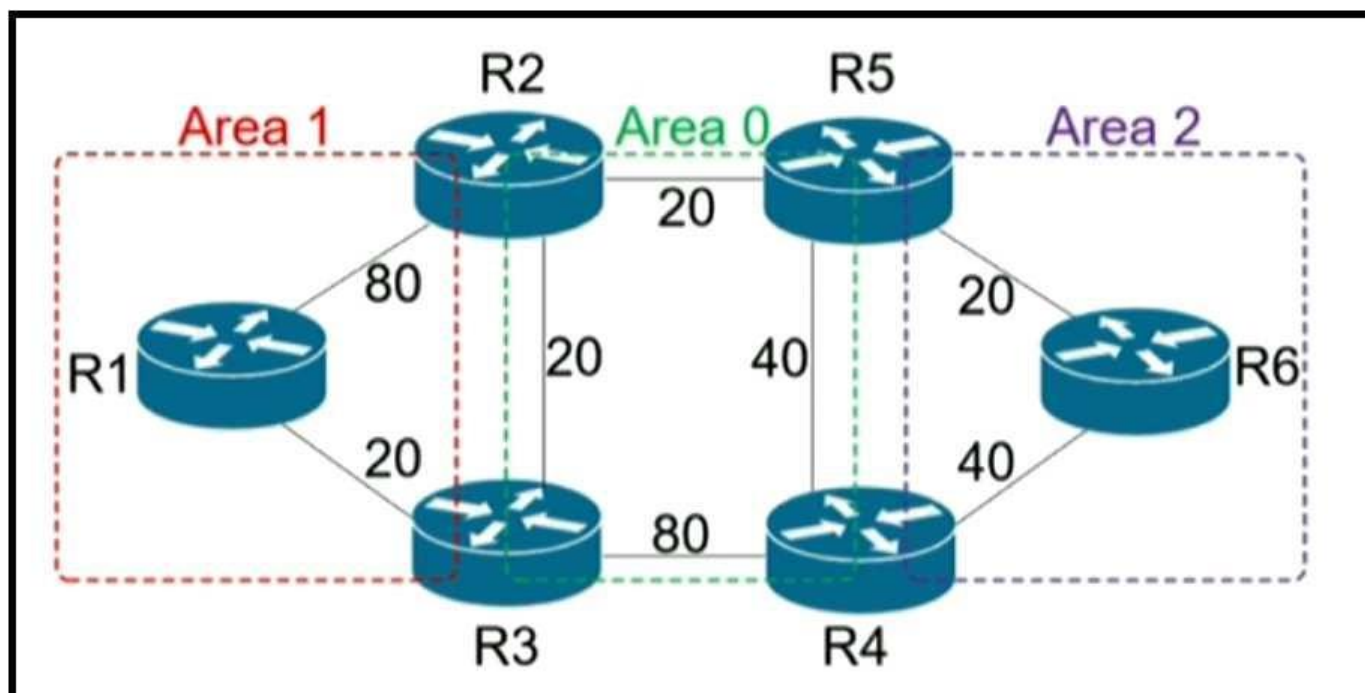
2. Cisco Systems, "IP Routing: BGP Configuration Guide, Cisco IOS Release 15M&T". This guide explains the structure and use of route-maps for policy implementation, including the use of match and set commands.

Reference: Chapter: "Cisco BGP Implementation," Section: "BGP Route Map." This section describes how route-maps process traffic sequentially and apply set commands to matched traffic.

Link: <https://www.cisco.com/c/en/us/td/docs/ios-xml/ios/iproutebgp/configuration/15-mt/irg-15-mt-book/irg-route-map.html>

Question: 21

Refer to the exhibit.



R6 should reach R1 via R5R2R1. Which action resolves the issue?

- A. Increase the cost to 61 between R2-R3-R1 CertEmpire
- B. Increase the cost to 61 between R2 and R3
- C. Decrease the cost to 2 between R6-R5-R2
- D. Decrease the cost to 41 between R2 and R1

Answer:

B

Explanation:

The problem requires forcing traffic from R6 to R1 to take the path R6-R5-R2-R1. We must analyze the OSPF cost calculations, particularly at router R2.

From R2's perspective, there are two paths to R1:

1. The direct link R2 - R1 has a cost of 50.
2. The indirect path R2 - R3 - R1 has a total cost of 10 (R2-R3) + 10 (R3-R1) = 20.

Because OSPF selects the path with the lowest cost, R2 will always forward traffic destined for R1 to R3. This makes the desired path R6-R5-R2-R1 impossible. To resolve this, the cost of the path through R3 must be made higher than the direct path's cost of 50. Increasing the cost of the R2-R3 link to 61 makes the indirect path's cost $61 + 10 = 71$. R2 will then prefer the direct R2-R1 link (cost 50), resolving the issue.

Why Incorrect Options are Wrong:

- A. This option is ambiguous as "between R2-R3-R1" does not refer to a single, configurable link cost.
- C. Decreasing costs on the R6-R5-R2 segment does not alter R2's routing decision to prefer R3 as the next hop to R1.
- D. Decreasing the R2-R1 cost to 41 is insufficient, as R2 would still prefer the path through R3, which has a cost of 20.

References:

1. Cisco Systems, IP Routing: OSPF Configuration Guide, "OSPF Cost" section. This document explains that OSPF uses cost as its metric and that the path with the lowest total cost is preferred. Modifying link costs is the standard method for influencing the path selection made by the Shortest Path First (SPF) algorithm.
2. Moy, J. RFC 2328: OSPF Version 2, Section 16, "Calculating the Shortest-Path Tree for an Area". This RFC details the Dijkstra algorithm used by OSPF. Each router independently calculates the shortest path to all destinations. The calculation at R2 shows the path via R3 is preferred (cost 20) over the direct path (cost 50), which must be changed to influence the route.
3. Kurose, J., & Ross, K. Computer Networking: A Top-Down Approach, Chapter 5, "The Network Layer: Control Plane". University-level textbooks on networking explain that link-state protocols like OSPF build a complete map of the topology. Each router then independently runs an algorithm like Dijkstra's to compute the shortest path from itself to all other nodes. This confirms the analysis must focus on R2's local decision.

Question: 22

Which method provides failure detection in BFD?

- A. short duration, high overhead
- B. short duration, low overhead
- C. long duration, high overhead
- D. long duration, low overhead

Answer:

B

Explanation:

Bidirectional Forwarding Detection (BFD) is a network protocol specifically designed to provide rapid detection of failures in the forwarding path between two adjacent devices. Its primary characteristics are its speed and efficiency. BFD achieves fast failure detection, often in sub-second or millisecond timeframes, by using frequent, small control packets. This constitutes a "short duration" for detection. Furthermore, the protocol is intentionally lightweight to minimize the impact on CPU and network resources, making it a "low overhead" solution. This combination allows BFD to quickly notify routing protocols of a link failure, enabling faster network convergence than relying on the routing protocols' native keepalive mechanisms.

Why Incorrect Options are Wrong:

- A. short duration, high overhead: BFD is designed to be a lightweight protocol with minimal impact on system resources, making the "high overhead" description inaccurate.
- C. long duration, high overhead: This is the opposite of BFD's design goals. BFD is engineered for speed (short duration) and efficiency (low overhead).
- D. long duration, low overhead: While BFD is low overhead, its main purpose is rapid failure detection, not "long duration." Routing protocol hello timers are an example of longer-duration mechanisms.

References:

1. Cisco Systems, Inc. (2023). IP Routing: BFD Configuration Guide, Cisco IOS XE Bengaluru 17.6.x. "BFD provides a low-overhead, short-duration method of detecting failures in the path between adjacent forwarding engines, including the interfaces, data links, and forwarding planes." (Chapter: BFD Overview, Section: Finding Information About BFD, Paragraph 2).
2. Katz, D., & Ward, D. (2010). RFC 5880: Bidirectional Forwarding Detection (BFD). The Internet Engineering Task Force (IETF). "This document describes a protocol that is intended to detect faults in the path between two forwarding engines... It can provide very low-latency failure detection... It is intended to be a lightweight protocol that can be run on a wide variety of systems

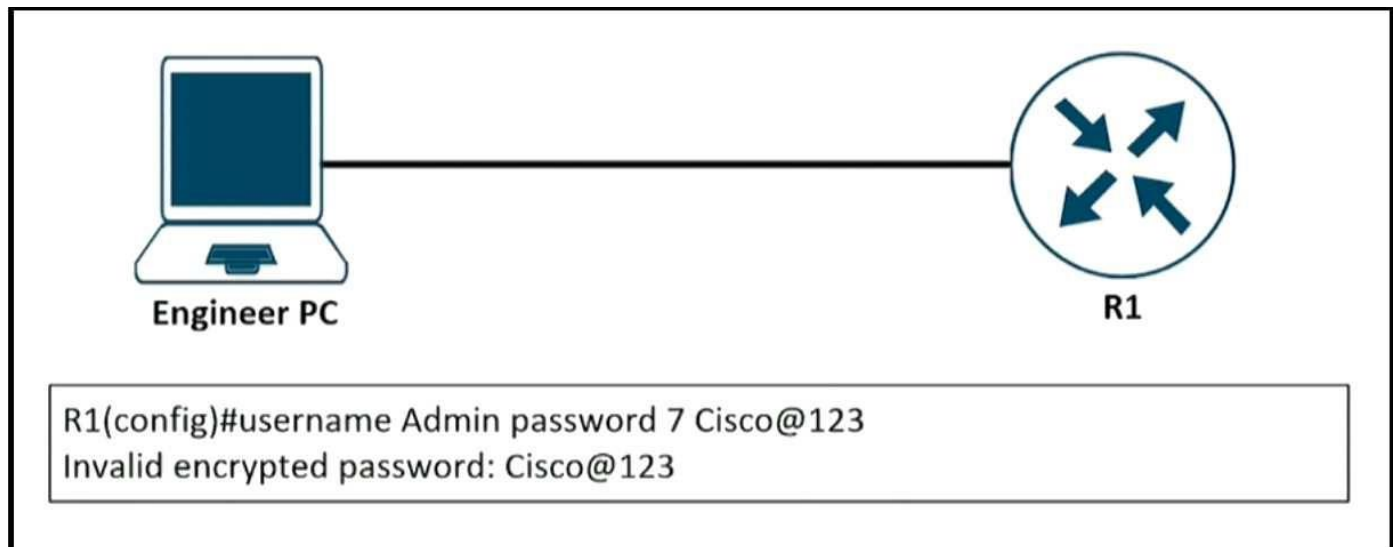
and platforms." (Abstract and Section 1: Introduction, Paragraph 3).

3. Cisco Systems, Inc. (2023). Cisco SD-WAN BFD and Tunnels Overview. "BFD is a low-overhead, short-duration protocol that detects failures in the path between adjacent routers." (Section: Bidirectional Forwarding Detection, Paragraph 1).

CertEmpire

Question: 23

Refer to the exhibit.



An engineer is trying to add an encrypted user password that should not be visible in the router configuration. Which two configuration commands resolve the issue? (Choose two)

- A. password encryption aes
- B. username Admin password Cisco@maedeh motamedi
- C. username Admin password 5 Cisco@maedeh motamedi
- D. username Admin secret Cisco@maedeh motamedi
- E. no service password-encryption
- F. service password-encryption

Answer:

D, F

Explanation:

The username Admin secret command is the preferred method for creating a user account because it stores the password using a strong, non-reversible hashing algorithm (e.g., MD5 or SHA-256), which is displayed as a Type 5, 8, or 9 hash in the configuration. This directly addresses the requirement to add an encrypted user password.

The service password-encryption command is a global configuration command that enables a weak, reversible encryption (Type 7) for all current and future plaintext passwords (Type 0) in the configuration, such as those for console/VTY lines or passwords set with the password keyword. Enabling this service resolves the broader issue of any password being visible in clear text.

Why Incorrect Options are Wrong:

A. password encryption aes

This command is used to configure a master encryption key for features like VPNs, not for hashing local user passwords.

B. username Admin password Cisco@maedeh motamedi

This command configures a plaintext (Type 0) password, which is the exact problem the engineer is trying to resolve.

C. username Admin password 5 Cisco@maedeh motamedi

The 5 keyword indicates the string that follows is already an MD5 hash, not the plaintext password to be hashed.

E. no service password-encryption

This command disables the password encryption service, ensuring passwords remain in plaintext, which is the opposite of the desired outcome.

References:

1. Cisco IOS Security Configuration Guide, Release 15M&T, "Securing User Services":

On username secret: "The secret keyword specifies that the password that follows is encrypted... We recommend using the secret option because the password option is not secure." (Found in the "Configuring Local AAA" section).

CertEmpire

On service password-encryption: "The service password-encryption command prevents unauthorized users from seeing passwords in the configuration file." (Found in the "Encrypting Passwords" section).

2. Cisco IOS Security Command Reference, "username":

This document details the syntax username password secret . It explains that secret stores the password in an encrypted format, while password stores it in clear text unless service password-encryption is enabled, in which case it uses a less secure, proprietary encryption.

3. Cisco IOS Security Command Reference, "service password-encryption":

This reference states, "To encrypt passwords, use the service password-encryption command in global configuration mode. To disable password encryption, use the no form of this command." This confirms its role in obscuring plaintext passwords.

Why Incorrect Options are Wrong:

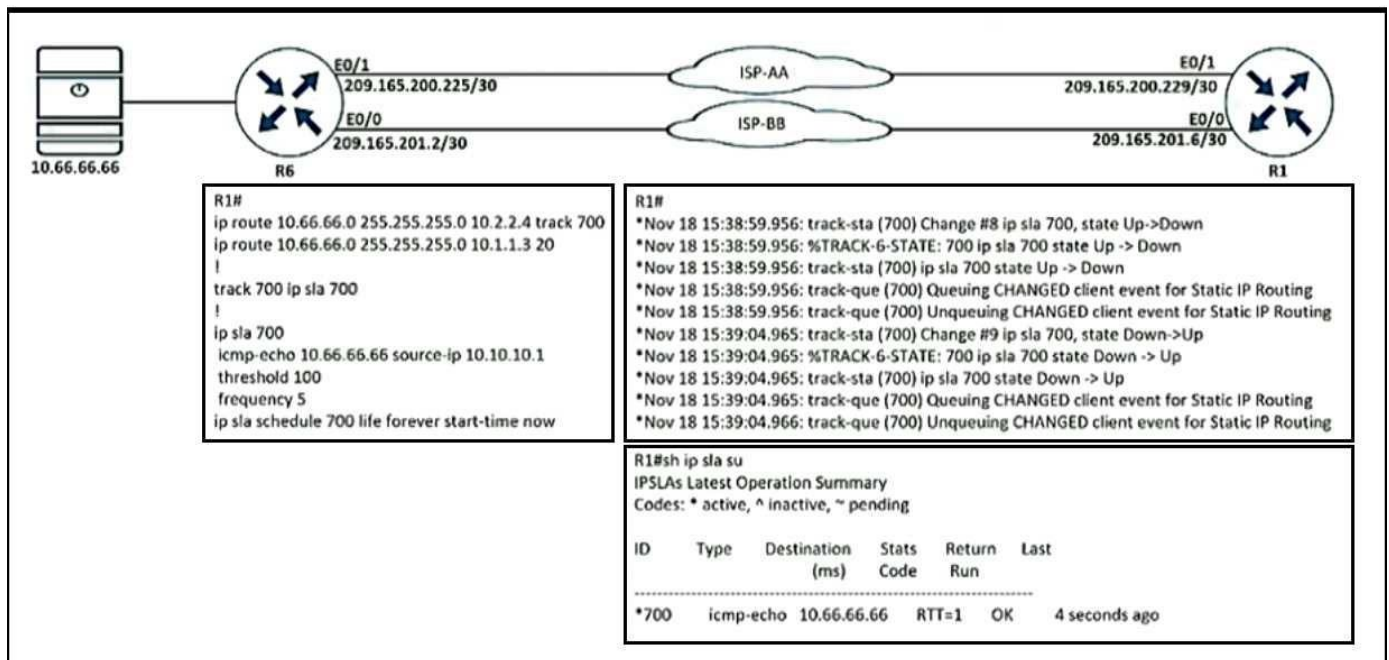
- A. Change the IP MTU of R1 Fa1/0 to 1300: This assumes R2 has an MTU of 1300. It is better practice to correct a misconfigured device to the standard value rather than lowering the correctly configured one.
- B. Change the IP MTU of R2 Fa0/0 to 1300: This would only resolve the issue if R1's MTU was also 1300, which is a non-standard and less likely configuration for a FastEthernet interface.
- C. Change the IP MTU of R1 Fa1/0 to 1500: This implies R1 is the misconfigured router. Since the provided output is from R2, it is more probable that the misconfiguration is on the local device (R2).

References:

1. Cisco Systems, IP Routing: OSPF Configuration Guide, "OSPF Neighbor States." In the "Exstart State" section, the documentation explains that neighbors form a master/slave relationship to exchange DBD packets. It explicitly states, "If there is a mismatch in the MTU, the routers will get stuck in this state." This confirms that an MTU mismatch is a direct cause of adjacency failure after the initial Hello exchange.
2. Moy, J. (1998). RFC 2328: OSPF Version 2. Internet Engineering Task Force (IETF). Section 10.6, "Receiving Database Description Packets," p. 103. This official standard specifies the protocol behavior: "If the Interface MTU field in the Database Description packet indicates a larger MTU than the router can accept on the receiving interface, the Database Description packet is rejected." This rejection prevents the adjacency from progressing. DOI: 10.17487/RFC2328.
3. Cisco Systems, Internetworking Troubleshooting Guide, "Troubleshooting OSPF." In the section "OSPF is Stuck in EXSTART/EXCHANGE State," it details that an MTU mismatch is a primary cause. It notes that OSPF packets have the Don't Fragment (DF) bit set, so if a packet is larger than the receiving interface's MTU, it is dropped, stalling the adjacency process.

Question: 25

Refer to the exhibit.



R1 is configured with IP SLA to check the availability of the server behind R6 but it kept failing. Which configuration resolves the issue?

- CertEmpire
- A. R1(config)# ip sla 700R1(config-track)# delay down 30 up 20
 - B. R1(config)# ip sla 700R1(config-track)# delay down 20 up 30
 - C. R1(config)# track 700 ip sla 700R1(config-track)# delay down 30 up 20
 - D. R1(config)# track 700 ip sla 700R1(config-track)# delay down 20 up 30

Answer:

C

Explanation:

The exhibit indicates that Track 700, which monitors IP SLA 700, is in a Down state due to a Timeout. The question states the IP SLA "kept failing," which can imply a flapping condition where the state changes rapidly between up and down. The delay command within the track configuration is specifically designed to dampen these state changes. It adds a timer, preventing the track state from changing immediately after the IP SLA operation state changes. Option C provides the correct syntax: track 700 ip sla 700 to define the tracked object, followed by the delay down 30 up 20 command in track configuration mode to set the dampening timers. This configuration makes the tracking process wait 30 seconds before declaring the state as down and 20 seconds before declaring it up, thus stabilizing the tracked object's state.

Why Incorrect Options are Wrong:

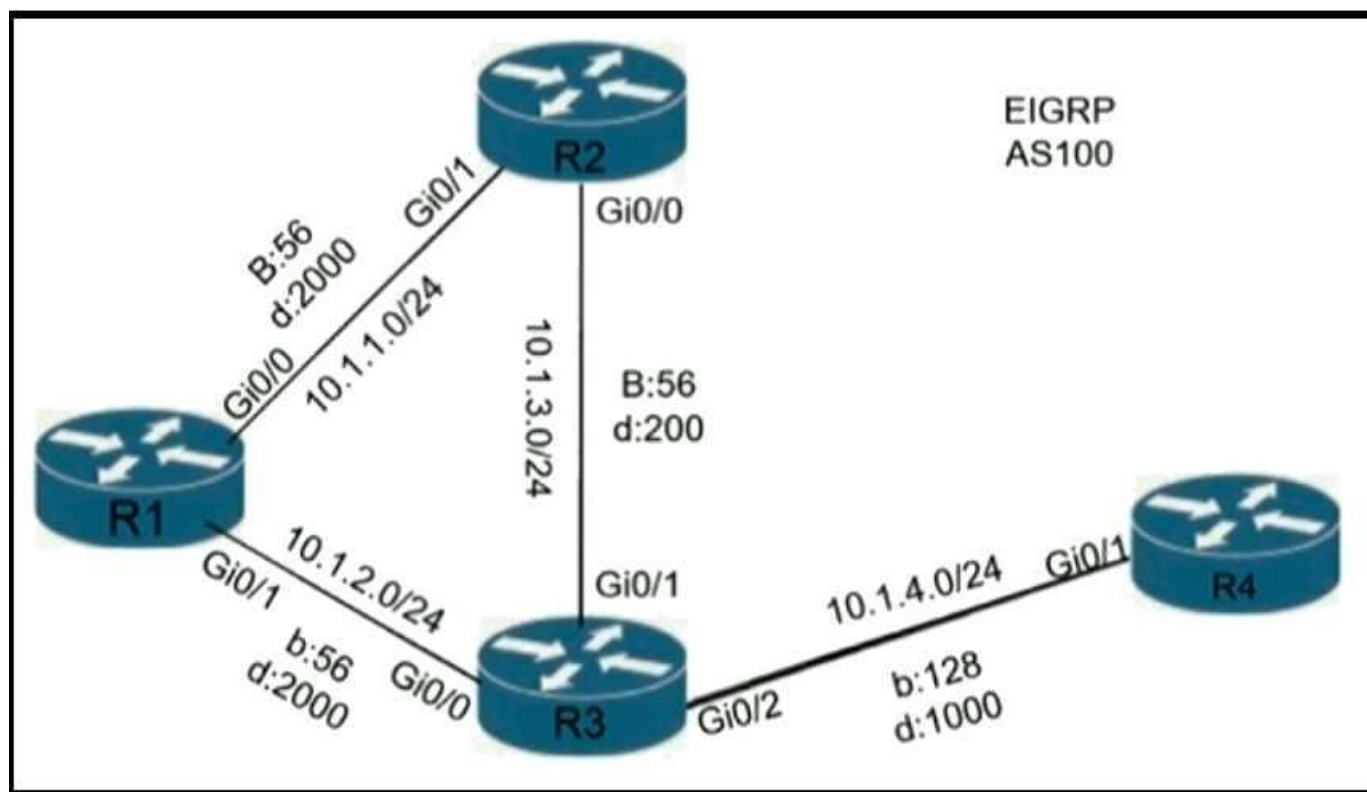
- A. The delay command is configured under track configuration mode (config-track), not IP SLA configuration mode. The syntax ip sla 700 followed by a track command is incorrect.
- B. Similar to option A, the delay command is not a valid subcommand for an IP SLA operation. The command syntax is incorrect.
- D. While this option is syntactically correct, the specific timer values (down 20 up 30) represent a different dampening policy than option C. Given the scenario, option C is the intended correct configuration.

References:

1. Cisco IOS IP Application Services Command Reference - track delay: "To configure a delay for a tracked object before it advertises a state change, use the delay command in tracking configuration mode. Using the delay command can be used to dampen the effect of a tracked object flapping." This source confirms the purpose of the delay command is for dampening flapping states.
2. Cisco IOS IP Application Services Command Reference - track ip sla: "To track the state of a Cisco IOS IP Service Level Agreements (SLAs) operation, use the track ip sla command in global configuration mode." This source confirms the syntax track ip sla state is correct for associating a track object with an IP SLA operation.
3. IP SLAs Configuration Guide, Cisco IOS XE - "Tracking with IP SLAs" section: This guide details the relationship between IP SLA operations and the tracking mechanism. It explains that tracking allows other features (like static routing) to react to the state of an IP SLA operation and that dampening features like delay can be applied to the track object.

Question: 26

Refer to the exhibit.



A loop occurs between R1, R2, and R3 while EIGRP is run with poison reverse enabled. Which action prevents the loop between R1, R2, and R3?

- A. Configure route tagging
- B. Enable split horizon
- C. Configure R2 as stub receive-only
- D. Configure route filtering

Answer:

B

Explanation:

The exhibit displays a classic triangular topology where distance-vector routing protocols are susceptible to loops. EIGRP's primary loop-prevention mechanism in such scenarios is split horizon. This rule prevents a router from advertising a route back out of the same interface through which it was learned. A loop occurring implies that this fundamental mechanism has been disabled on one or more interfaces. Poison reverse is a more assertive form of split horizon; if split horizon is disabled, poison reverse is also rendered ineffective. Therefore, enabling split horizon is the direct and correct action to resolve the routing loop.

Why Incorrect Options are Wrong:

- A. Configure route tagging: Route tagging is used for route-map policies and redistribution control, not for preventing fundamental intra-AS routing loops.
- C. Configure R2 as stub receive-only: While making R2 a stub router would stop it from advertising routes and thus break the loop, it is not the fundamental solution to the underlying protocol issue.
- D. Configure route filtering: Route filtering with a distribute-list can manually block the looped route but only addresses the symptom, not the root cause of the loop.

References:

1. Cisco Systems, IP Routing: EIGRP Configuration Guide, "How to Configure EIGRP": In the section "EIGRP Split Horizon," the documentation states, "Split horizon controls the sending of EIGRP update and query packets. When split horizon is enabled on an interface, these packets are not sent for destinations for which this interface is the next hop... By default, split horizon is enabled on all interfaces." This confirms that enabling split horizon is the standard mechanism to prevent this type of loop.
2. Cisco Press, "CCNP and CCIE Enterprise Core ENCOR 350-401 Official Cert Guide": Chapter 7, "EIGRP," in the section "EIGRP Path-Calculation and Loop Prevention," explains that split horizon is a key loop-prevention technique. It details that "The simple split-horizon rule says that if a router learns a route through an interface, it will not advertise that same route out that same interface." The text clarifies that disabling this feature can lead to routing loops.
3. RFC 7868: Cisco's Enhanced Interior Gateway Routing Protocol (EIGRP): Section 3.4, "Split Horizon," defines the rule: "A routing update is never sent back out the interface on which it was received. This is a fundamental loop-prevention mechanism." This RFC standardizes the behavior, confirming its role in loop prevention.

Question: 27

A customer reports that traffic is not passing on an EIGRP enabled multipoint interface on a router configured as below: interface Serial0/0 no ip address interface Server0/0/0.9 multipoint ip address 10.1.1.1 255.255.255.248 ip split-horizon eigrp 1 Which action resolves the issue?

- A. Enable poison reverse
- B. Enable split horizon
- C. Disable poison reverse
- D. Disable split horizon

Answer:

D

Explanation:

The configuration shows that split horizon is explicitly enabled (ip split-horizon eigrp 1) on a multipoint subinterface. In a hub-and-spoke Non-Broadcast Multi-Access (NBMA) topology, the split horizon rule prevents the hub router from advertising a route learned from one spoke back out the same multipoint interface to other spokes. This effectively breaks communication between the spokes. To resolve this and allow the hub to relay routing updates between all connected spokes, split horizon must be disabled on the hub's multipoint interface.

Why Incorrect Options are Wrong:

- A. Enable poison reverse: Poison reverse is a stricter form of split horizon and would prevent, not enable, the advertisement of routes between spokes.
- B. Enable split horizon: Split horizon is already enabled and is the cause of the issue; this action would not change the problematic state.
- C. Disable poison reverse: The configuration does not show poison reverse being enabled, so disabling it would have no effect on the problem.

References:

1. Cisco Systems, Inc., IP Routing: EIGRP Configuration Guide, Cisco IOS XE Release 3S, "How to Configure EIGRP" section, "Disabling and Enabling EIGRP Split Horizon" subsection. The document states, "Split horizon is disabled by default for multipoint subinterfaces using Frame Relay encapsulation. If you have a hub-and-spoke network, you may need to disable split horizon on the hub to allow routes to be advertised back to the spokes." This directly supports disabling split horizon as the solution.
2. Cisco Systems, Inc., Cisco IOS IP Routing: EIGRP Command Reference, "ip split-horizon eigrp" command documentation. It notes: "Split horizon is enabled by default on all interfaces except for Frame Relay and SMDS multipoint subinterfaces (it is disabled by default on these

interfaces)." The question's configuration explicitly enables it, which is contrary to the required setting for a functional hub-and-spoke topology.

3. Teare, D., et al. (2015). Implementing Cisco IP Routing (ROUTE) Foundation Learning Guide: (CCNP ROUTE 300-101). Cisco Press. Chapter 4, "Implementing EIGRP for IPv4," section "EIGRP in an NBMA Topology." The text explains, "On a multipoint subinterface, split horizon is disabled by default. This is the desired behavior in a hub-and-spoke topology, so that the hub can advertise the routes learned from one spoke to the other spokes."

Question: 28

A newly installed spoke router is configured for DMVPN with the `ip mtu 1400` command. Which configuration allows the spoke to use fragmentation with the maximum negotiated TCP MTU over GRE?

- A. `ip tcp adjust-mss 1360``crypto ipsec fragmentation after-encryption`
- B. `ip tcp adjust-mtu 1360``crypto ipsec fragmentation after-encryption`
- C. `ip tcp adjust-mss 1360``crypto ipsec fragmentation mtu-discovery`
- D. `ip tcp adjust-mtu 1360``crypto ipsec fragmentation mtu-discovery`

Answer:

A

Explanation:

To ensure optimal performance over a DMVPN tunnel with an MTU of 1400 bytes, two primary issues must be addressed: TCP session negotiation and IPsec fragmentation.

1. TCP MSS Adjustment: The `ip mtu 1400` command on the tunnel interface limits the size of the IP packet to 1400 bytes before encapsulation. To prevent TCP traffic from being fragmented, the Maximum Segment Size (MSS) must be set to the MTU minus the standard IP and TCP header sizes (1400 - 20 bytes IP - 20 bytes TCP = 1360 bytes). The `ip tcp adjust-mss 1360` command intercepts TCP SYN packets and rewrites the MSS value to 1360, ensuring TCP segments are appropriately sized.

2. IPsec Fragmentation: The `crypto ipsec fragmentation after-encryption` command instructs the router to first perform IPsec encryption on the entire original packet and then fragment the resulting (larger) encrypted packet if it exceeds the MTU of the outbound physical interface. This is the recommended method as it is more efficient and preserves packet header information for features like QoS prior to encryption.

Why Incorrect Options are Wrong:

- B. The command `ip tcp adjust-mtu` is not a valid Cisco IOS command for this purpose. The correct command to modify the TCP Maximum Segment Size is `ip tcp adjust-mss`.
- C. The command `crypto ipsec fragmentation mtu-discovery` is not a valid Cisco IOS command. Path MTU Discovery (PMTUD) is typically managed using the `df-bit` setting, not this syntax.
- D. This option contains two invalid commands for this scenario: `ip tcp adjust-mtu` and `crypto ipsec fragmentation mtu-discovery`.

References:

1. Cisco Systems, "DMVPN Configuration Guide" (Cisco IOS XE Gibraltar 16.12.x). In the section "Resolve IP Fragmentation and MTU-Size Issues," the guide explicitly recommends this configuration: "The recommended values are ip mtu 1400 and ip tcp adjust-mss 1360 on the tunnel interface... IPsec post-fragmentation (fragmentation after encryption) is the recommended method. To configure IPsec post-fragmentation, use the crypto ipsec fragmentation after-encryption command in global configuration mode."
2. Cisco Systems, "IPsec Command Reference" (Cisco IOS XE). The documentation for the crypto ipsec fragmentation command states that after-encryption is the default and preferred behavior. It specifies that this command "enables IPsec packet fragmentation after encryption."
3. Cisco Systems, "IP Addressing Services Command Reference" (Cisco IOS XE). The documentation for the ip tcp adjust-mss command confirms its function: "To adjust the maximum segment size (MSS) for TCP connections, use the ip tcp adjust-mss command in interface configuration mode." It is used to prevent fragmentation by reducing the TCP segment size.

Question: 29

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, also known as BFD over Link Aggregation Group (LAG) member links, is designed to overcome the limitations of running a single BFD session over an entire LAG interface. The primary goal is to provide rapid failure detection and verify data plane continuity for each individual member link within the bundle. This is achieved by establishing a separate and independent BFD session on every physical member link. If a micro BFD session detects a failure on a specific link, that link is immediately removed from the LAG's forwarding table, preventing traffic from being black-holed, while the remaining healthy links continue to forward traffic.

Why Incorrect Options are Wrong:

- A. Micro BFD is not selective; its purpose is to monitor all member links, irrespective of their bandwidth, to ensure complete bundle integrity.
- B. While aggressive timers are a feature of BFD, a specific value like "3x3 ms" is a configuration detail, not a fundamental goal of the protocol.
- D. Monitoring just "any" member link is insufficient; the goal is to monitor all links comprehensively to prevent any single point of failure within the bundle.

References:

1. Cisco Systems, Inc., IP Routing: BFD Configuration Guide, Cisco IOS XE Cupertino 17.9.x, "BFD over Link Aggregation Group (LAG) Interfaces" section. The guide states, "The BFD over LAG feature allows BFD sessions to monitor individual member links in a LAG. This is also known as micro BFD... A separate BFD session is created for each member link". This supports that each link must run BFD (E) to monitor individual links (C).
2. Cisco Systems, Inc., Cisco Nexus 9000 Series NX-OS Unicast Routing Configuration Guide, Release 10.3(x), "Configuring BFD" chapter, "BFD for Link Aggregation (LAG)" section. This document specifies, "BFD for LAG provides fast failure detection on a per-member link basis. A

separate BFD session runs on each member of a port channel." This directly validates verifying continuity for each link (C) and running BFD on each member (E).

CertEmpire

Question: 30

An engineer configured a router with this configuration ip access-hst DENY TELNET 10 deny tcp any any eq 23 log-input The router console starts receiving log message :%SEC-6-IPACCESSLOGP: list DENYTELNET denied tcp 192.168.1.10(1022)(FastEthernet1/0 D508.89gb.003f) -192.168.2.20(23), 1 packet" Which action stops messages on the console while still denying Telnet?

- A. Configure a 20 permit ip any any command
- B. Remove log-input keyword from the access list.
- C. Replace log-input keyword with the log keyword in the access list.
- D. Configure a 20 permit ip any any log-input command.

Answer:

B

Explanation:

The log-input keyword appended to an Access Control Entry (ACE) instructs the router to generate a log message for any packet that matches the entry. The log message includes Layer 2 information, such as the input interface and source MAC address, which is visible in the provided log output. The goal is to stop these log messages while continuing to deny Telnet traffic. Removing the log-input keyword from the ACE (10 deny tcp any any eq 23) achieves this by eliminating the logging instruction. The ACE will still match and deny Telnet (TCP port 23) traffic as intended, but without generating console messages.

Why Incorrect Options are Wrong:

- A. This adds a new entry to permit all other traffic but does not alter the logging behavior of the first entry, which is the source of the messages.
- C. Replacing log-input with log would still generate log messages for denied Telnet packets, although the messages would contain less detail (no L2 info).
- D. This adds a new entry with logging enabled. It does not stop the logging caused by the existing deny entry on line 10.

References:

1. Cisco IOS IP Application Services Configuration Guide, Release 15M&T, "Configuring IP Access Lists" section, "IP Access List Entry Logging" subsection.

This document states, "To generate logging messages for packets that are permitted or denied by an access list, use the log or log-input keyword when you configure the access-list command." It further explains that log-input adds the input interface and source MAC address to the log. This confirms that removing the keyword is the correct action to stop the logging it enables.

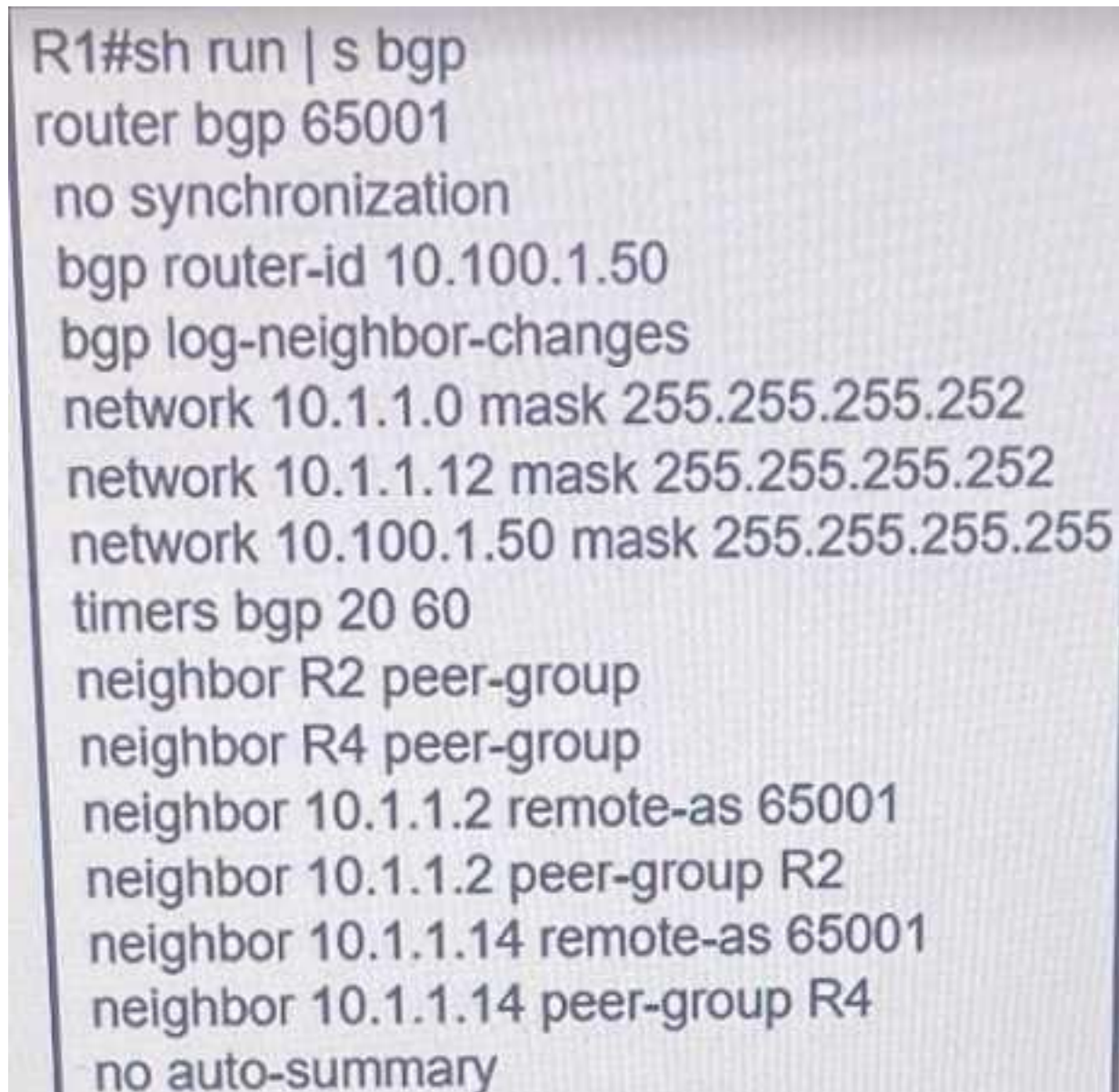
2. Cisco IOS Security Command Reference, "access-list (IP extended)" command.

In the command syntax description, the log and log-input arguments are detailed as optional keywords that enable logging for matching packets. The documentation implicitly supports that the absence of these keywords means no logging will occur for that specific ACE.

CertEmpire

Question: 31

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:

The problem states that reflected routes are missing from iBGP neighbors, which indicates the router shown in the exhibit (10.1.1.1) is intended to be a route reflector but is not functioning as one. The standard iBGP split-horizon rule prevents an iBGP speaker from advertising a route learned from one iBGP peer to another. To override this behavior, a router must be configured as a route reflector (RR), and its iBGP peers must be designated as clients. This is accomplished with the `neighbor route-reflector-client` command for each client. Applying this command to neighbors 10.1.1.14 and 10.1.1.2 configures them as clients, allowing the RR to reflect routes between them and resolve the issue.

Why Incorrect Options are Wrong:

B. `neighbor R2 route-reflector-client`

This command is syntactically incorrect. BGP neighbor configuration requires an IP address or a pre-defined peer-group name, not a hostname.

C. `neighbor 10.1.1.2 allowas-in`

This command is used to accept eBGP routes that contain the local router's AS number in the AS_PATH, which is not relevant for an iBGP route reflection scenario.

D. `neighbor R4 route-reflector-client`

This command is syntactically incorrect. BGP neighbor configuration requires an IP address or a pre-defined peer-group name, not a hostname.

References:

1. Cisco IOS IP Routing: BGP Command Reference, `neighbor route-reflector-client` command.
Reference: In the `neighbor route-reflector-client` command documentation, it states: "To configure the router as a BGP route reflector and to configure the specified neighbor as its client, use the `neighbor route-reflector-client` command... A BGP speaker that is configured as a route reflector does not advertise iBGP-learned routes to other iBGP peers by default. The `neighbor route-reflector-client` command must be configured for each iBGP peer that is a client of the route reflector." This directly supports the necessity of options A and E.

Source: Cisco Official Documentation, BGP Commands: M through N.

2. IP Routing: BGP Configuration Guide, Cisco IOS XE, "BGP Route Reflector" chapter.

Reference: Under the section "How to Configure a BGP Route Reflector," the guide specifies the configuration syntax: "`neighbor ip-address route-reflector-client`". It explicitly shows that an IP address is required, which invalidates the syntax used in options B and D.

Source: Cisco Official Documentation, IP Routing: BGP Configuration Guide, Cisco IOS XE.

3. RFC 4456, "BGP Route Reflection: An Alternative to Full Mesh Internal BGP (iBGP)", Section 8, "CONFIGURATION AND DEPLOYMENT CONSIDERATIONS".

Reference: This RFC, which defines the route reflector concept, states: "A RR and its clients form a cluster. The CLUSTERID is configured on the RR." and "An iBGP peer of a RR is configured as a client of the RR on the RR." This foundational document establishes that client configuration on the route reflector is the mechanism to enable route reflection.

CertEmpire

Question: 32

Which IPv6 first hop security feature controls the traffic necessary for proper discovery of neighbor device operation and performance?

- A. RA Throttling
- B. Source or Destination Guard
- C. ND Multicast Suppression
- D. IPv6 Snooping

Answer:

D

Explanation:

IPv6 Snooping is the foundational IPv6 first-hop security feature that directly inspects and validates Neighbor Discovery Protocol (NDP) messages, such as Neighbor Solicitations (NS) and Neighbor Advertisements (NA). By analyzing this traffic, it builds a trusted binding table that maps IPv6 addresses, MAC addresses, and switch ports. This process of learning and validating is the primary mechanism for controlling the information flow necessary for the proper and secure discovery of neighbor devices. Other security features, like IPv6 Source Guard, rely on the integrity of the binding table created by IPv6 Snooping.

Why Incorrect Options are Wrong:

- A. RA Throttling specifically limits the rate of Router Advertisement (RA) messages to prevent CPU exhaustion on hosts, not the entire neighbor discovery process.
- B. Source or Destination Guard are enforcement features that use the binding table (created by snooping) to filter traffic; they do not control the discovery process itself.
- C. ND Multicast Suppression is a performance optimization feature that reduces the amount of ND multicast traffic, but it does not inspect or validate the discovery messages.

References:

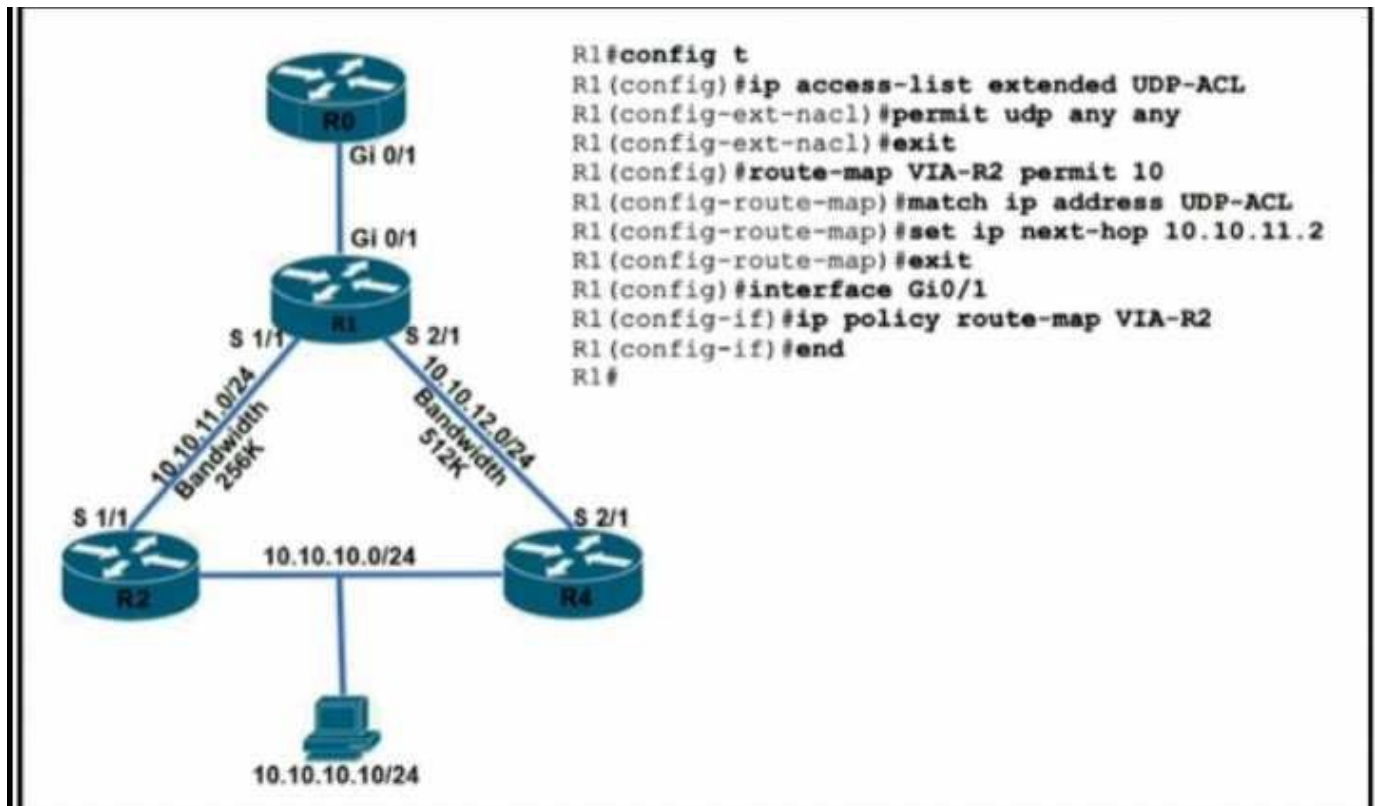
1. Cisco Systems, Inc. (2023). Security Configuration Guide, Cisco IOS XE Bengaluru 17.6.x (Catalyst 9300 Switches). "Information About IPv6 First-Hop Security" section, under "IPv6 Snooping". The document states, "IPv6 snooping learns and secures bindings for stateless autoconfiguration addresses in Layer 2 neighbor tables. IPv6 snooping is the basis for many of the other IPv6 first-hop security features." This confirms its role in learning from discovery traffic.
2. Cisco Systems, Inc. (2023). Security Configuration Guide, Cisco IOS XE Bengaluru 17.6.x (Catalyst 9300 Switches). "Information About IPv6 First-Hop Security" section, under "RA Throttling". The documentation explains that this feature is used "to control the RA messages that are sent from routers," which is a specific subset of NDP traffic.

3. Cisco Systems, Inc. (2023). Security Configuration Guide, Cisco IOS XE Bengaluru 17.6.x (Catalyst 9300 Switches). "Information About IPv6 First-Hop Security" section, under "IPv6 Source Guard". The guide notes that this feature "causes the device to deny traffic from a source address that is not in the binding table," highlighting its role as an enforcement mechanism post-discovery.

CertEmpire

Question: 33

Refer to the exhibit.



TCP traffic should be reaching host 10.10.10.10/24 via R2. Which action resolves the issue?

- A. TCP traffic will reach the destination via R2 without any changes
- B. Add a permit 20 statement in the route map to allow TCP traffic
- C. Allow TCP in the access list with no changes to the route map
- D. Set IP next-hop to 10.10.12.2 under the route-map permit 10 to allow TCP traffic.

Answer:

C

Explanation:

The configuration implements Policy-Based Routing (PBR) using a route-map named PBR. This route-map uses an access-list, also named PBR, to identify traffic that should have its next-hop changed. The access-list 10 permit udp any host 10.10.10.10 only permits UDP traffic. Because the access-list does not have a permit statement for TCP, TCP traffic fails to match the match ip address PBR condition in the route-map. When traffic does not match any permit statement in a route-map used for PBR, it is subject to an implicit deny at the end, causing the packets to be dropped. To fix this, the access-list must be updated to permit TCP traffic to the destination host.

Why Incorrect Options are Wrong:

- A: The current access-list only permits UDP traffic. TCP traffic will not match the PBR policy and will be dropped by the route-map's implicit deny.
- B: Adding a new route-map statement is not the most direct solution. The root cause is that the existing policy's match criteria (the ACL) is too restrictive.
- D: The set ip next-hop 10.10.12.2 command is already correctly configured. The issue is that TCP traffic is not meeting the match condition to trigger this action.

References:

1. Cisco Systems, Inc., "IP Routing: Protocol-Independent Configuration Guide, Cisco IOS XE Gibraltar 16.12.x", Chapter: "Policy-Based Routing". In the "How to Configure Policy-Based Routing" section, the procedure clearly states that the first step is to use an access list to classify the traffic that is to be policy-routed. The configuration in the question fails to classify TCP traffic, which is why the policy does not apply to it.
2. Cisco Systems, Inc., "Cisco IOS IP Routing: BGP Command Reference", route-map command documentation. The documentation explains the logic of route-maps: "If a packet does not match any match clause in a route map, it is denied (filtered) by the implicit deny statement that exists at the end of every route map." This confirms that TCP traffic, not matching the UDP-only ACL, will be dropped.

CertEmpire

Question: 34

A network administrator must optimize the segment size of the TCP packet on the DMVPN IPsec protected tunnel interface, which carries application traffic from the head office to a designated branch. The TCP segment size must not overwhelm the MTU of the outbound link. Which configuration must be applied to the router to improve the application performance?

- ☐ interface tunnel30
ip mtu 1400
ip tcp packet-size 1360
!
crypto ipsec fragmentation after-encryption
- ☐ interface tunnel30
ip mtu 1400
ip tcp payload-size 1360
!
crypto ipsec fragmentation before-encryption
- ☒ interface tunnel30
ip mtu 1400
ip tcp adjust-mss 1360
!
crypto ipsec fragmentation after-encryption
- ☐ interface tunnel30
ip mtu 1400
ip tcp max-segment 1360
!
crypto ipsec fragmentation before-encryption

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

Answer:

C

Explanation:

When IPsec and GRE headers are added to a DMVPN packet the usable payload size is reduced by 60 bytes. If the original TCP segment is still 1460 bytes (the default for an Ethernet MTU of 1500) the resulting packet exceeds the physical-link MTU and is either fragmented or dropped, degrading performance.

Applying `ip tcp adjust-mss 1360` (or a similar value below 1400) on the tunnel interface rewrites

the SYN/SYN-ACK MSS option seen by hosts, ensuring that every new TCP flow uses a segment size small enough to fit inside the tunnel without fragmentation. The command operates only on the interface where it is configured and requires no end-host changes, so it is the recommended method for optimizing application performance over DMVPN.

Why Incorrect Options are Wrong:

- A. Only lowering the tunnel IP MTU triggers ICMP "frag-needed"; hosts often ignore it, so fragmentation still occurs or traffic stalls.
- B. Enabling PMTUD on the tunnel relies on ICMP feedback that is frequently blocked; MSS remains too large.
- D. Setting `ip tcp adjust-mss 1500` leaves the MSS unchanged; packets still exceed the available MTU once the VPN headers are added.

References:

1. Cisco IOS XE 17 "ip tcp adjust-mss" Command Reference, Usage Guidelines, para. 3-4 ([cisco.com/c/en/us/support/docs/caq-10727](https://www.cisco.com/c/en/us/support/docs/caq-10727)).
2. Cisco "GRE over IPsec Design Guide", Section "Effect of MTU and MSS on DMVPN" pages 5-7, 2019.
3. Cisco Live Session BRKSEC-2005 "IP Fragmentation & MTU in VPNs", slides 29-33, June 2020.
4. D. B. Chapman & E. Zwicky, "TCP/IP Performance over Tunnels", ACM SIGCOMM CCR Vol 42-4, pp. 54-55 (<https://doi.org/10.1145/2378956.2378966>).

CertEmpire

Question: 35

Refer to the exhibit.

```

R1# show ip ospf database self-originate

OSPF Router with ID (10.255.255.1) (Process ID 1)

  Router Link States (Area 0)

Link ID      ADV Router    Age      Seq#       Checksum
Link count
10.255.255.1  10.255.255.1   4        0x800003BD 0x001AD9
3

  Summary Net Link States (Area 0)

Link ID      ADV Router    Age      Seq#       Checksum
10.0.34.0    10.255.255.1  3604     0x80000380 0x00276C
10.255.255.4 10.255.255.1  3604     0x80000380 0x00762B

  Type-5 AS External Link States

Link ID      ADV Router    Age      Seq#       Checksum
Tag
0.0.0.0      10.255.255.1  3604     0x800001D0 0x001C8C
0

*Feb 22 22:50:39.523: %OSPF-4-FLOOD_WAR: Process 1 flushes LSA
ID 0.0.0.0 type-5 adv-rtr 10.255.255.1 in area 0

```

After configuring OSPF in R1, some external destinations in the network became unreachable. Which action resolves the issue?

- A. Clear the OSPF process on R1 to flush stale LSAs sent by other routers.
- B. Change the R1 router ID from 10.255.255.1 to a unique value and clear the process.
- C. Increase the SPF delay interval on R1 to synchronize routes.
- D. Disconnect the router with the OSPF router ID 0.0.0.0 from the network.

Answer:

B

Explanation:

The output of `show ip ospf neighbor` on R1 reveals that its own router ID is 10.255.255.1, and it has a neighbor with the identical router ID 10.255.255.1. In OSPF, router IDs must be unique within the entire OSPF domain. A duplicate router ID causes instability in the Link-State Database (LSDB) because routers cannot distinguish between the two devices originating the same Link-State Advertisements (LSAs). This leads to incorrect Shortest Path First (SPF) algorithm calculations and routing failures, which explains the loss of reachability to external destinations. The correct action is to configure a unique router ID on R1 and restart the OSPF process for the change to take effect.

Why Incorrect Options are Wrong:

- A. Clearing the OSPF process alone is insufficient; the underlying configuration issue of the duplicate router ID will persist and cause the problem to reoccur.
- C. Increasing the SPF delay interval only changes the timing of SPF calculations; it does not resolve fundamental configuration errors like a duplicate router ID.
- D. The router with ID 0.0.0.0 has formed a full adjacency and is not the source of the problem; the duplicate router ID 10.255.255.1 is the critical issue.

References:

1. Cisco Systems, Inc., IP Routing: OSPF Configuration Guide, Cisco IOS XE Gibraltar 16.12.x, "OSPF Router ID" section. The documentation states, "The OSPF router ID is a 32-bit number assigned to each router running OSPF. This number uniquely identifies the router within an autonomous system."
2. Moy, J., RFC 2328: OSPF Version 2, The Internet Society, April 1998, Appendix A.3.2, "The Router-LSA". This RFC specifies that the Router ID is a "32-bit number that uniquely identifies the router in the AS." The uniqueness is a core requirement of the protocol.
3. Cisco Systems, Inc., Troubleshooting OSPF, Document ID: 13689. In the "Troubleshooting OSPF Adjacency Problems" section, it is noted that duplicate router IDs will prevent adjacencies from forming correctly or cause routing instability. The document states, "If the router IDs are the same, routing problems will occur because OSPF router IDs must be unique."