

Cisco Certified Design Expert Practical Exam

Cisco 352-011

Version Demo

Total Demo Questions: 33

Total Premium Questions: 336

Buy Premium PDF

<https://passqueen.com>

support@passqueen.com

QUESTION NO: 1 - (DRAG DROP)

DRAG DROP

When developing a multicast network design, SSM should be used for which type of source and receiver distribution?

limited sources	Source Distribution
many sources	Target
limited receivers	Receiver Distribution
many receivers	Target

ANSWER:

When developing a multicast network design, SSM should be used for which type of source and receiver distribution?

limited sources	Source Distribution
many sources	limited sources
limited receivers	Receiver Distribution
many receivers	many receivers

Explanation:

Source-Specific Multicast is designed around receivers explicitly identifying both the multicast group and the desired source, commonly expressed as an (S,G) channel. This model works best when the multicast content originates from a limited and identifiable set of sources, such as a small collection of video servers, conferencing bridges, market-data publishers, or other controlled application publishers. Because the source is known by receivers, the network can build multicast distribution state specifically for that source and group, rather than requiring receivers to receive traffic from every possible sender to a group.

The receiver side of this design can scale to many receivers. Multiple endpoints can subscribe to the same known source and multicast group, and the network replicates traffic only along the paths needed to reach those interested listeners. This is the normal one-to-many distribution model that makes multicast efficient for high-bandwidth content. A limited source population also makes source discovery, authorization, operational control, and receiver configuration more predictable. Cisco describes SSM as a multicast model in which receivers request traffic from particular sources, using IGMPv3 membership reporting. See Cisco's [Source-Specific Multicast configuration documentation](#) and the [IETF SSM service model](#). Accordingly, the correct distribution pattern is limited sources with many receivers.

QUESTION NO: 2

What are two possible drawbacks of ending Loop-Free Alternate to support fast convergence for most destination IGP prefixes? (Choose two)

- A. The IGP topology might need to be adjust
- B. Loop-free alternate's convergence in less than 100 milliseconds is not possible
- C. Loop-free alternate's are supported only for prefixes that are considered external tot the IGP
- D. Loop-free alternates are not supported in global VPN VRF OSPF instances
- E. Additional path computations are needed

ANSWER: A E

Explanation:

Enabling Loop-Free Alternate (LFA) protection can require that the IGP topology be adjusted because a router needs an eligible neighboring router that satisfies the loop-free condition for each protected destination. In some physical topologies, especially hub-and-spoke or sparsely meshed designs, no neighbor meets that condition for all prefixes. Network designers may therefore need to add links, alter link metrics, or otherwise improve path diversity to obtain broad LFA coverage. This is a design tradeoff: the topology and metric plan must support both normal shortest-path forwarding and a precomputed, loop-free backup next hop.

Additional path computations are also needed because the router must evaluate candidate neighbors and determine whether each one qualifies as a loop-free alternate for relevant destinations. This work is performed in addition to the normal shortest-path calculations used by the IGP. The result is valuable fast local repair, since the backup can be installed before a failure occurs, but it introduces extra computational and control-plane planning considerations as the number of prefixes, neighbors, and protection requirements grows. Cisco's LFA documentation describes the precomputed backup-path model, while [RFC 5286](#) defines the loop-free alternate criteria and coverage limitations. See also Cisco's [Loop-Free Alternate configuration guide](#).

QUESTION NO: 3

What is a design aspect regarding multicast transport for MPLS Layer 3 VPNs using the Rosen Draft implementation?

- A. LDP is the multicast control plane protocol.
- B. Multicast traffic is forwarded over GRE tunnels.
- C. Multicast traffic is forwarded over LDP or RSVP signaled LSPs.
- D. Using the MDT SAFI in BGP ensures that PIM can be disabled in the core.

ANSWER: B**Explanation:**

Multicast traffic is forwarded over GRE tunnels. The Rosen Draft multicast VPN model builds multicast distribution trees between provider-edge routers so that a multicast source in one VPN can reach interested receivers at other provider-edge routers while preserving VPN separation. The ingress provider-edge router encapsulates customer multicast packets in Generic Routing Encapsulation, with an outer provider multicast header used to carry the packet across the service-provider network. The provider network then replicates and forwards this encapsulated traffic along a default multicast distribution tree, or along a data multicast distribution tree when traffic volume warrants a more selective tree. This architecture allows multicast forwarding state in the core to represent the provider multicast transport tree rather than requiring customer multicast routes to be installed throughout the core. In design terms, the GRE encapsulation and multicast distribution-tree behavior are defining characteristics of the original Rosen approach to multicast MPLS Layer 3 VPNs. Cisco's multicast VPN configuration documentation describes Rosen-based MVPN operation and its default and data MDT concepts: [Cisco IOS Multicast VPN Configuration Guide](#). The underlying Rosen specification also defines encapsulation of VPN multicast packets for transport over provider multicast distribution trees: [RFC 6037](#).

QUESTION NO: 4

Which DCI technology utilizes a "flood and learn" technique to populate the Layer 2 forwarding table?

- A. OTV
- B. E-VPN
- C. VPLS
- D. LISP

ANSWER: A

Explanation:

OTV is correct because it extends Layer 2 connectivity between geographically separate data centers by creating an overlay across an intervening Layer 3 transport network. OTV performs MAC-address learning in the data plane: when a frame sourced by an unknown MAC address arrives, the OTV edge device learns that source MAC address and associates it with the appropriate overlay interface. For traffic sent to an unknown destination MAC address, OTV uses controlled flooding across the overlay, subject to its multicast and adjacency mechanisms. This flood-and-learn behavior allows the OTV edge devices to build their Layer 2 forwarding information dynamically without requiring endpoint MAC addresses to be preconfigured. OTV also uses an authoritative edge device to ensure that broadcast, unknown-unicast, and multicast traffic is injected into the overlay only when appropriate, helping prevent unnecessary duplication between sites. Cisco describes OTV as a Layer 2 overlay technology for data-center interconnect that preserves Ethernet forwarding semantics while using an IP-based transport core. See Cisco's [Overlay Transport Virtualization overview](#) and the [Cisco OTV configuration documentation](#).

QUESTION NO: 5

Refer to the Exhibit.

In which three Layers should you use nonstop Forwarding to reduce service impact in case of failure? (Choose three)

- A. Enterprise Edge
- B. Enterprise Core
- C. Service provider Edge
- D. Dual-attached sever Layer
- E. Enterprise Non-Routed Access Layer
- F. Enterprise Non-Routed Distribution Layer.

ANSWER: A B C

Explanation:

Nonstop Forwarding is most valuable at routing-system boundaries and in the routed backbone, where a control-plane switchover must not unnecessarily interrupt forwarding of traffic that is already supported by hardware forwarding tables. **Enterprise Edge, Enterprise Core, and Service provider Edge** are routed infrastructure layers that carry transit traffic and commonly form critical failure domains. Enabling NSF in these layers allows the forwarding plane to continue using previously programmed forwarding information while the routing control plane recovers or a redundant route processor takes over. This reduces packet loss and preserves service continuity during planned maintenance or certain supervisor, route-processor, and protocol-process failures.

NSF depends on resilient forwarding hardware and, for many routing protocols, NSF-aware neighbors that can maintain adjacency and continue forwarding during the recovery interval. It is therefore a component of a broader high-availability design that also includes redundant devices, redundant links, fast failure detection, and appropriate routing convergence behavior. Cisco describes NSF as a high-availability capability intended to maintain packet forwarding across a routing-engine switchover; see [Cisco IOS XE NSF-aware routing documentation](#) and [Cisco Nonstop Forwarding overview](#).

QUESTION NO: 6

A network design engineer is designing a new storage area network that combines multiple separate legacy SAN environments within a data center. Which technology isolates events within one of the SAN environment from the others?

- A. FCIP tunnels
- B. N-port ID Virtualization

C. N-Port Virtualization

D. Virtual SANs

ANSWER: D

Explanation:

Virtual SANs are correct because a VSAN provides a logically separate Fibre Channel fabric over shared physical SAN switching infrastructure. Each VSAN has its own fabric services and maintains an independent control-plane and data-plane environment, including its own zoning database, name server, fabric login processes, and domain-related fabric operations. Consequently, disruptive fabric events—such as a reset, reconfiguration, or state-change notification—are contained within the affected VSAN rather than propagating to devices in the other VSANs. This makes VSANs particularly appropriate for consolidating separate legacy SAN fabrics in one data center while preserving operational isolation, administrative separation, and fault containment. Cisco MDS switches support multiple isolated VSANs on the same physical switch and links; an inter-VSAN routing service is required only when controlled communication between VSANs is explicitly needed. This architecture enables infrastructure consolidation without turning distinct SAN environments into a single shared failure or event domain. Cisco documents VSANs as virtual fabrics that provide isolation among devices and traffic, with separate instances of fabric services for each VSAN. See the [Cisco MDS NX-OS VSAN configuration guide](#) and the [Cisco MDS SAN virtualization overview](#).

QUESTION NO: 7

Which two options are Loop-Free Alternate design considerations? (Choose two)

- A. MPLS TE must be enabled because it is used for building the backup paths
- B. Backup coverage and effectiveness is dependent on the network topology
- C. It can simplify the capacity planning by matching the backup path with the post-convergence path
- D. It provides an optional backup path by avoiding low bandwidth and edge links
- E. It can impact SLA-sensitive appliance by routing traffic to low bandwidth links while IGP convergence is in progress

ANSWER: B E

Explanation:

Backup coverage and effectiveness is dependent on the network topology because a Loop-Free Alternate (LFA) is a precomputed neighbor-based repair path that must satisfy the loop-free condition for the protected destination. Whether a router has a qualifying alternate next hop is therefore determined by the physical and logical topology, link metrics, and placement of routers. A design should evaluate LFA coverage for each relevant failure scenario rather than assume that every prefix and adjacency can be protected. RFC 5286 defines the loop-free condition and explains that the availability of alternate paths depends on topology and routing distances.

It can impact SLA-sensitive appliance by routing traffic to low bandwidth links while IGP convergence is in progress because LFA provides immediate local repair after a failure, before normal IGP reconvergence selects the final post-failure shortest path. The available loop-free neighbor can be reachable through a path with less capacity or different performance characteristics than the normal path. Network designers must consequently validate bandwidth, delay, and loss behavior of repair paths for critical traffic, particularly where links are asymmetric or capacity constrained. Cisco documents LFA as an IP fast-reroute mechanism using precomputed backup next hops. See [RFC 5286](#) and [Cisco IP Routing: Loop-Free Alternate Configuration Guide](#).

QUESTION NO: 8

When is it required to leak routes into an IS-IS Level 1 area?

- A. when equal cost load balancing is required between the backbone and non-backbone areas

- B. when unequal cost load balancing is required between the backbone and non-backbone areas
- C. when MPLS L3VPN PE devices are configured in the Level 1 areas
- D. when a multicast RP is configured in the non-backbone area

ANSWER: C

Explanation:

Route leaking is required when MPLS L3VPN PE devices are configured in the Level 1 areas. In a normal IS-IS hierarchy, a Level 1 router learns a default route toward the nearest Level 1/Level 2 router for destinations outside its area, rather than learning the individual Level 2 prefixes. That default-route behavior is usually sufficient for ordinary IP forwarding, but MPLS L3VPN PE routers need reachability to remote PE loopback addresses through the MPLS transport core. Those PE loopback prefixes must be available as specific IGP routes so that MPLS label distribution and label-switched forwarding can establish the required end-to-end transport path between PE devices.

Leaking the appropriate Level 2 reachability information, commonly remote PE loopback /32 prefixes, into the Level 1 area provides that specific routing knowledge while retaining the IS-IS two-level design. Cisco documents IS-IS route leaking as the mechanism used to advertise selected prefixes between Level 1 and Level 2 domains, enabling connectivity requirements that cannot rely solely on the Level 1 default route. See Cisco's [IS-IS Route Leaking Configuration Guide](#) and its [MPLS Layer 3 VPN Configuration Guide](#).

QUESTION NO: 9

When designing fast convergence on a network using loop-free alternate, on which two basis can the next-hop routes be precomputed? (Choose two)

- A. Per neighbor
- B. Per network type
- C. Per link
- D. Per prefix
- E. Per failure type

ANSWER: C D

Explanation:

Loop-Free Alternate provides local IP fast reroute by identifying a loop-free backup next hop before a failure occurs. The backup information can be calculated and installed on a **per-link** basis or on a **per-prefix** basis. Per-link protection precomputes an alternate that protects traffic when the primary outgoing link fails. This is useful where protecting the directly connected forwarding path is the primary design goal. Per-prefix protection evaluates the alternate path for each individual destination prefix, allowing the router to select a backup next hop that is loop-free for that specific destination. This offers more granular protection because alternate-path eligibility can differ by destination, even when the same primary next hop is used.

Both approaches rely on the LFA loop-free condition: the selected neighbor must be able to reach the destination without sending the traffic back through the repairing router after the protected failure. The router performs these calculations from link-state topology information and places eligible repair paths into the forwarding plane so traffic can be redirected immediately, without waiting for normal routing-protocol convergence. RFC 5286 defines the fundamental LFA mechanisms and repair-path calculation model: [RFC 5286](#). Cisco's OSPF fast-reroute documentation provides implementation context for LFA-based local repair: [Cisco IOS XE OSPF LFA documentation](#).

QUESTION NO: 10

A data center provider has designed a network using these requirements

Two data center sites are connected to the public internet

Both data centers are connected to different Internet providers

Both data centers are also directly connected with a private connection for the internal traffic can also be at this direct connection The data center provider has only /19 public IP address block

Under normal conditions, Internet traffic should be routed directly to the data center where the services are located. When one Internet connections fails to complete traffic for both data centers should be routed by using the remaining Internet connection in which two ways can this routing be achieved? (Choose two)

- A.** One /20 block is used for the first data center and the second /20 block is used for the second data center. The /20 block from the local data center is sent out without path prepending and the /20 block from the remote data center is sent out with path prepending at both sites
- B.** One /20 block is used for the first data center and the second /20 block is used for the second data center. Each /20 block is only sent out locally. The /19 block is sent out at both Internet connections for the backup case to reroute the traffic through the remaining internet connection
- C.** One /20 block is used for the first data center and the second /20 block is used for the second data center. The /20 block from the local data center is sent out with a low BGP local preference and the /20 block from the remote data center is sent out with a higher BGP local preference of both sites
- D.** BGP will always load-balance the traffic to both data center sites
- E.** One /20 block is used for the first data center and the second /20 block is used for the second data center. The /20 block from the local data center is sent out with a low BGP weight and the /20 block from the remote data center is sent out with a higher BGP weight at both sites
- F.** The data center provider must have an additional public IP address block for this routing

ANSWER: A B

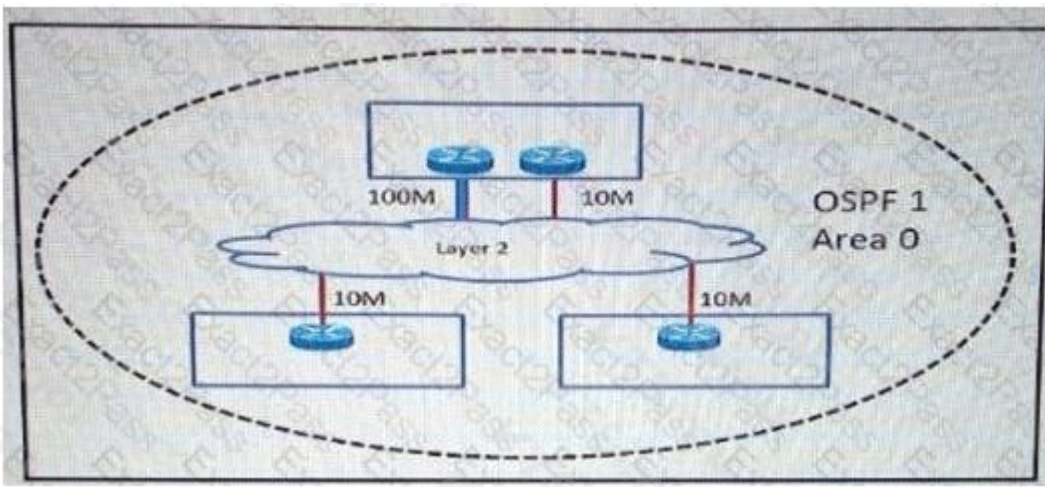
Explanation:

Using two /20 subnets from the provider's /19 allocation enables each data center to advertise the more-specific prefix for the services it hosts. Advertising each local /20 normally and advertising the remote /20 with AS-path prepending makes the locally originated path more attractive to external networks under normal conditions. If an Internet connection fails, the surviving site continues advertising reachability for both prefixes and uses the private inter-data-center connection to deliver traffic to services at the other location. AS-path length is a standard BGP path-selection attribute, so prepending is a commonly used inbound-traffic-engineering technique when the provider cannot directly control remote networks' routing policies.

The aggregate-and-more-specific design also achieves the required active/backup behavior. Each site advertises only its local /20 during normal operation, while both sites advertise the covering /19 aggregate. External networks select the /20 because BGP forwarding uses longest-prefix matching. Upon loss of a site's Internet connectivity, its /20 advertisement is withdrawn; traffic then matches the /19 still advertised through the surviving Internet connection. The private connection carries that traffic between data centers as necessary. Cisco describes BGP best-path processing at [Cisco BGP Best Path Selection Algorithm](#); the BGP base specification is [RFC 4271](#).

QUESTION NO: 11

Refer to the exhibit.



An enterprise has three sites over a Layer 2 Metro Ethernet ELAN service. 100Mb/s and 10 Mb/s links have been provisioned to provide redundancy for the head office. When OSPF routing enabled to provide connectivity and the correct bandwidth statement has been applied to each interface, the branch sites observe two equal-cost routes to the head office. The enterprise wants to send all traffic through the 100 Mb/s link and use the 10Mb/S link strictly as a backup. Which OSPF network type must be set to ensure that the head office 100 Mb/s circuit is preferred over the 10 Mb/s circuit, at the same time minimize the amount of configuration required on all of the routers throughout the network?

- A. NBMA
- B. Point-to-multipoint
- C. Point-to-point
- D. Broadcast

ANSWER: C

Explanation:

Point-to-point is the appropriate OSPF network type. Although Metro Ethernet ELAN is a multipoint Layer 2 service, configuring OSPF point-to-point on the participating Ethernet interfaces avoids the broadcast-network DR/BDR pseudonode behavior that can hide the far-end circuit cost when SPF calculates a path across the shared segment. OSPF instead represents the connectivity as router-to-router point-to-point links, allowing the cumulative OSPF metric to reflect the configured costs associated with the 100 Mb/s and 10 Mb/s head-office circuits. With interface bandwidth correctly set, the 100 Mb/s circuit receives the lower OSPF cost and is selected as the primary path; the 10 Mb/s circuit remains available as a higher-cost backup.

This approach also avoids the additional neighbor statements and manual adjacency maintenance commonly required by nonbroadcast designs. Point-to-point operation has no DR or BDR election and uses normal OSPF neighbor discovery on Ethernet, keeping the deployment operationally simple while ensuring deterministic primary/backup path selection. Cisco's OSPF design guidance discusses how network type affects adjacencies and SPF topology representation: [Cisco OSPF Design Guide](#). The underlying OSPF point-to-point link and SPF metric behavior is specified in [RFC 2328](#).

QUESTION NO: 12

At which two networks points is route summarization supported? (Choose two)

- A. At EIGRP AS boundaries
- B. At EIGRP interface boundaries
- C. At OSPF virtual-link boundaries
- D. At EIGRP are boundaries
- E. At OSPF area boundaries

ANSWER: B E

Explanation:

Route summarization is supported at EIGRP interface boundaries and at OSPF area boundaries. In EIGRP, manual summarization is configured on the outbound interface with the `ip summary-address eigrp` command. The router then advertises the aggregate route through that specific interface instead of advertising the individual component prefixes. This provides a deliberate point of control for limiting routing-table size and containing the scope of topology changes.

In OSPF, interarea route summarization occurs at an Area Border Router, which is the boundary between an attached nonbackbone area and the backbone or another area. An ABR can advertise a summary for routes learned within an area into other areas, using the `area range` configuration. This hierarchy is fundamental to OSPF scalability because it reduces the number of interarea prefixes propagated beyond an area. Cisco documentation describes EIGRP summary-address configuration per interface and OSPF area-range summarization at ABRs: [Cisco EIGRP Route Summarization](#) and [Cisco OSPF Route Summarization](#).

QUESTION NO: 13

A network engineering team is in the process of designing a lab network for a customer demonstration. The design engineer wants to show that the resiliency of the MPLS traffic Engineering Fast Reroute solution has the same failover/failback times as a traditional SONET/SDH network (around 50MSEC). In order to address both link failure and node failure within the lab topology network, which type of the MPLS TE tunnels must be considered for this demonstration?

- A. TE backup tunnel
- B. Next-hop (NHop) tunnel
- C. FRR Backup tunnel
- D. next-next-hop (NNHop) tunnel

ANSWER: D

Explanation:

The correct choice is **next-next-hop (NNHop) tunnel**. MPLS Traffic Engineering Fast Reroute provides local protection by pre-establishing a backup path at the point of local repair, allowing traffic to be redirected immediately after a protected failure is detected. This local switchover is designed to provide sub-50-millisecond restoration, comparable to the protection objective traditionally associated with SONET/SDH networks.

An NNHop backup tunnel is specifically built to bypass both the directly connected next-hop router and the failed link leading to it. The point of local repair sends protected traffic to a downstream router beyond the next hop, so the traffic can continue even when either the adjacent TE link or the adjacent transit node becomes unavailable. This is the required protection model when the demonstration must validate resilience against both link and node failures rather than only a single local link failure. Cisco documents Fast Reroute next-next-hop protection as the mechanism used for node protection, with the backup path merging at a router downstream of the failed node. See Cisco's [MPLS TE Fast Reroute documentation](#) and the [Cisco IOS XE MPLS TE Fast Reroute guide](#).

QUESTION NO: 14

Your client is considering acquiring a new IPv6 address block so that all Ethernet interfaces on the network receive addresses based on their burned-in hardware addresses, with support for 600 VLANs. Which action do you recommend?

- A. Acquire a new /60 IPv6 network and subnet it into /70 networks, one per VLAN
- B. Acquire a new /58 IPv6 network and subnet it into /64 networks, one per VLAN
- C. Acquire a new /60 Ipv6 network and subnet it into /68 networks, one per VLAN

D. Acquire a new /54 IPv6 network and subnet it into /64 networks , one per VLAN

ANSWER: D

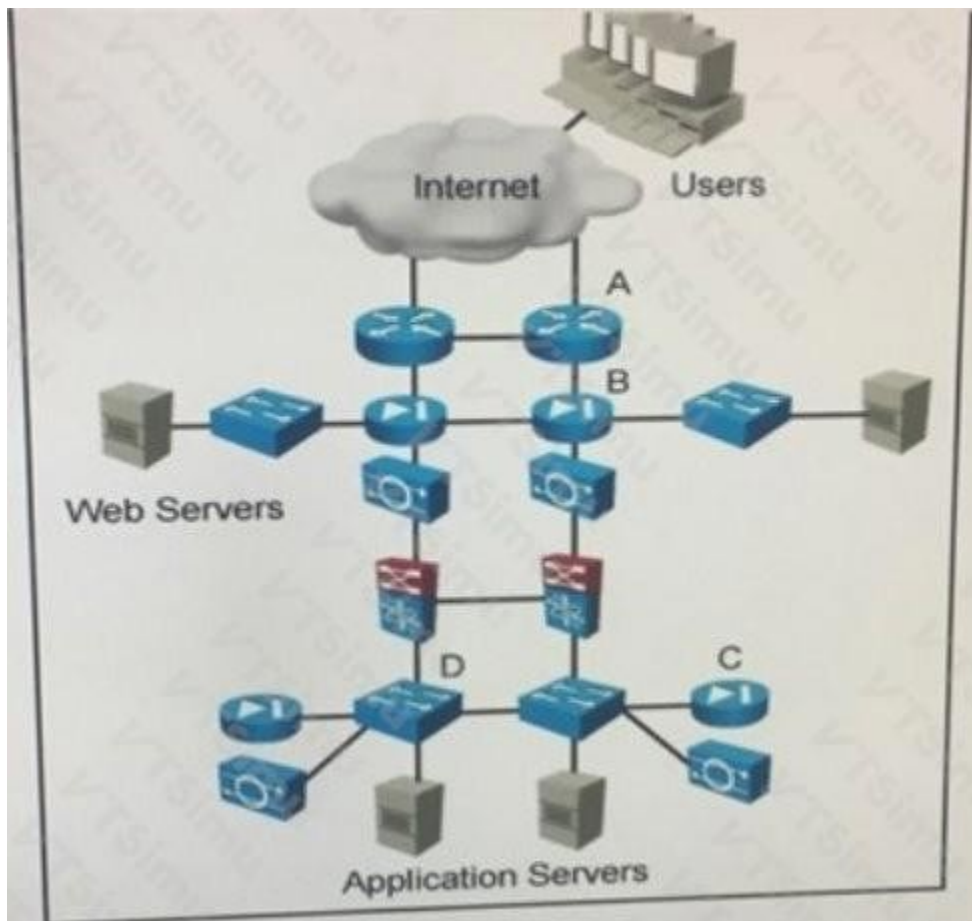
Explanation:

Acquire a new /54 IPv6 network and subnet it into /64 networks , one per VLAN. Ethernet interfaces using an interface identifier derived from a burned-in MAC address use the modified EUI-64 format, which occupies 64 bits of the IPv6 address. Consequently, each VLAN that requires this addressing model needs a /64 subnet: the first 64 bits identify the routed subnet and the remaining 64 bits hold the interface identifier. This is also the standard IPv6 subnet size used by Stateless Address Autoconfiguration and is the broadly applicable design choice for LAN segments.

A /54 allocation leaves 10 bits between the organizational prefix length and the /64 subnet boundary. Those 10 subnet bits provide 2^{10} , or 1,024, separate /64 VLAN prefixes. This exceeds the requirement for 600 VLANs while retaining a single contiguous, aggregatable allocation. The IPv6 addressing architecture defines the 64-bit interface-ID convention for prefixes used with this format, and RFC 7421 documents the operational rationale for the /64 boundary in IPv6 networks. See [RFC 4291, IPv6 Addressing Architecture](#) and [RFC 7421, Analysis of the 64-bit Boundary](#).

QUESTION NO: 15

Refer to the exhibit.



The operations team has identified that some of the multi-tiered e-commerce application have slow performance, due to illegitimate inbound traffic form the internet. On which device do you place traffic filtering to improve performance?

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

ANSWER: D

Explanation:

Option D is correct because it is the Internet-facing ingress enforcement point in the illustrated path. Filtering illegitimate inbound traffic at that point stops unwanted packets before they consume downstream capacity and processing resources used by the e-commerce application. This placement protects the available bandwidth on the Internet connection and prevents unnecessary traffic from reaching internal security devices, load-balancing functions, and application tiers. As a result, valid customer traffic has greater access to the resources required for normal application operation, directly addressing the observed performance degradation.

Cisco access-control design guidance recommends applying traffic controls as close to the traffic source as practical when doing so does not interfere with required policy or routing behavior. For untrusted Internet-originated traffic, the perimeter ingress point is therefore an appropriate location for filtering policies. The filter should be deliberately scoped to permit the public services the application requires while denying unauthorized, malformed, or otherwise illegitimate traffic. Cisco also describes edge-based protections as an important component of defending infrastructure and preserving service availability during unwanted traffic events. See Cisco's [access control list guidance](#) and [DDoS protection overview](#).

QUESTION NO: 16

Which four resources does Cisco Cloud Center provision in an ACL environment? (Choose four)

- A. VLAN Pool
- B. Contracts
- C. End point Group (EPG)
- D. VRF
- E. Subject/Filters
- F. Application Network Profile (ANP)

ANSWER: B C E F

Explanation:

Cisco CloudCenter's integration with Cisco Application Centric Infrastructure provisions the application-policy objects needed to connect an application deployment while enforcing its intended segmentation and communication policy. An Application Network Profile (ANP) provides the application-centric container for the deployed application's networking policy. Within that profile, an End point Group (EPG) identifies a logical group of application endpoints that share common network and security requirements.

Contracts express the allowed communications between endpoint groups. The detailed policy that a contract applies is represented through subjects and filters: a subject associates one or more filters with a contract, and filters define the traffic classification criteria, such as protocol and port information. Together, Application Network Profile (ANP), End point Group (EPG), Contracts, and Subject/Filters form the ACI policy constructs CloudCenter uses to model the application and its permitted connectivity.

Cisco's ACI documentation describes application profiles, EPGs, contracts, subjects, and filters as the core application-policy model: [Cisco ACI Developer Documentation](#). Cisco also documents CloudCenter as an application deployment and lifecycle-management platform: [Cisco CloudCenter](#).

QUESTION NO: 17

Which markup language is used to format Ansible 's playbook?

- A. ADML
- B. YAML

- C. HTML
- D. XML
- E. NAML

ANSWER: B

Explanation:

YAML is the language used to write and format Ansible playbooks. Ansible playbooks are human-readable automation files that define one or more plays, including the target hosts, tasks, variables, handlers, and configuration details needed to automate operations. YAML's indentation-based structure represents nested data clearly, which makes playbooks easy to read, version-control, and maintain. In a playbook, YAML mappings commonly define attributes such as `hosts`, `tasks`, and module parameters, while YAML lists are used for collections such as multiple tasks. Although YAML is often described as a data-serialization language rather than a markup language in a strict technical sense, it is the expected and correct format for Ansible playbooks. Ansible documentation specifies that playbooks are written in YAML and notes that correct indentation and YAML syntax are essential because they determine the structure interpreted by Ansible. See the [Ansible Playbooks introduction](#) and the [Ansible YAML syntax reference](#) for the official guidance.

QUESTION NO: 18

Which feature must be part of the network design to wait a predetermined amount of time before notifying the routing protocol of a change in the path in the network?

- A. Transmit delay
- B. Throttle timer
- C. SPF hold time
- D. Interface dampening

ANSWER: D

Explanation:

Interface dampening is the appropriate feature because it deliberately delays routing-protocol notification of an interface state transition for a configured interval. This controls the impact of transient link events, such as a physical interface briefly going down and returning to service, before those events are advertised into the routing domain. By withholding immediate notification, the design avoids unnecessary routing updates, repeated route withdrawals and installations, and avoidable control-plane convergence activity caused by unstable interfaces.

In Cisco IOS, IP Event Dampening is specifically intended to reduce the effects of interface flapping on routing protocols. The configured parameters govern how interface events are dampened and when the routing protocol is informed, making this feature directly aligned with the requirement to wait a predetermined amount of time before propagating a path-change event. It should be applied thoughtfully, with timers appropriate for the network's convergence and availability objectives, because intentional notification delay is a trade-off between stability and the speed at which a real failure is reflected in routing. Cisco documents the behavior and configuration of this capability in its [IP Event Dampening documentation](#) and provides related routing design guidance in the [Cisco routing convergence overview](#).

QUESTION NO: 19

Which two statements regarding to QoS marking are true? (Choose two)

- A. Shaping is one of the ways that packets can be remarked
- B. Class-based marking occurs after packet classification
- C. 802.1Q/p CoS bits and IP Precedence are both layer 3 marking fields

D. QoS marking establishes a trust boundary that scheduling tools depend on

E. MPLS EXP and DSCP are both layer 2 marking fields

ANSWER: B D

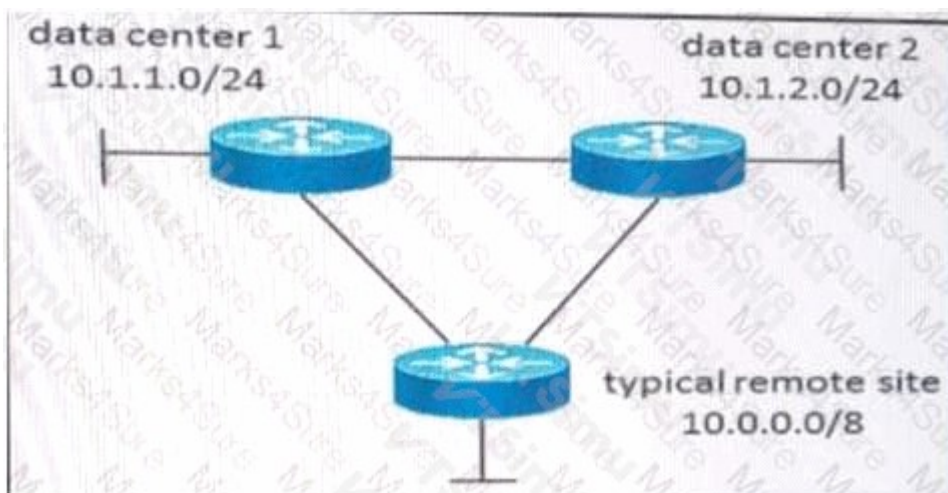
Explanation:

Class-based marking occurs after packet classification because the device must first identify the traffic class to which a packet belongs before it can apply or modify a QoS marking. Cisco MQC policies use classification criteria, such as access control lists, DSCP values, protocols, or input interfaces, and then apply actions—including setting DSCP, IP precedence, or Layer 2 CoS values—to packets in the matching class. This sequence lets a policy consistently assign markings according to the organization's QoS design.

QoS marking establishes a trust boundary that scheduling tools depend on because marking defines the QoS treatment a packet is intended to receive as it crosses the network. At a trust boundary, a device either accepts markings from a trusted source or classifies and rewrites markings from an untrusted source. Subsequent QoS functions, including queuing, congestion management, and scheduling, use those trusted classification and marking values to place traffic into the appropriate queues and deliver the designed service behavior. Cisco describes this relationship between classification, marking, trust boundaries, and later QoS actions in its [QoS classification documentation](#) and [QoS policy documentation](#).

QUESTION NO: 20

Refer to the exhibit.



A customer currently has a large EIGRP-based network with several remote sites attached. All remote sites connect to the two corporate data centers, depicted as 10.1.1.0 and 10.1.2.0. The customer has experienced several network-wide failures where neighbors were stuck-in-active and had other network stability issues due to some links flapping. Which two redesign options increase stability and reduce the load on the remote site routers, still maintaining optimal routing between remote sites and the two data centers? (Choose two)

A. Set the data center routers as stub-routers

B. Perform summarization at the data centers, selectively leaking routes sent to the remote sites

C. Perform summarization at the remote sites, selectively leaking routes sent to the data centers

D. Set the hello interval timer to be larger than the hold interval

E. Increase the hold interval to accommodate lost hello packets on error-prone links

ANSWER: B E

Explanation:

Perform summarization at the data centers, selectively leaking routes sent to the remote sites is appropriate because EIGRP queries can propagate broadly when a route is lost and no feasible successor exists. Advertising summaries toward the remote sites creates a query boundary at the data-center edge, reducing the number of specific prefixes that branch routers must maintain and process. Selective route leaking preserves visibility for any more-specific routes that are required for optimal forwarding, while the aggregate handles the remaining destination space. This approach reduces control-plane work on remote routers and contains the impact of topology changes.

Increase the hold interval to accommodate lost hello packets on error-prone links improves neighbor stability when transient loss affects hello packets. EIGRP declares a neighbor down only after the hold time expires without receiving a qualifying hello or other EIGRP packet. A suitably increased hold time prevents brief errors or intermittent loss from repeatedly tearing down a valid adjacency. Avoiding unnecessary adjacency resets also avoids the resulting route withdrawals, recalculations, and query activity that can contribute to stuck-in-active events. The hold time must be engineered together with the hello interval so failures are still detected within the business-required convergence time. Cisco EIGRP behavior and timer processing are documented in [Cisco's EIGRP documentation](#); the protocol's route-query and convergence mechanisms are specified in [RFC 7868](#).

QUESTION NO: 21

Which two options describe the advantages of using DWDM over traditional optical networks? (Choose two)

- A. Inherent topology flexibility with intelligent chromatic dispersion
- B. Inherent topology flexibility and service protection provided without penalty through intelligent oversubscription of bandwidth reservation
- C. Inherent topology flexibility with built-in service protection
- D. Inherent topology flexibility with a service protection provided through a direct integration with an upper layer protocol
- E. Ability to expand bandwidth over existing optical infrastructure

ANSWER: A E

Explanation:

DWDM increases the capacity and operational flexibility of an optical transport network by carrying many independently modulated wavelengths over one fiber pair. "Ability to expand bandwidth over existing optical infrastructure" is therefore a fundamental benefit: operators can add wavelengths and raise aggregate capacity without deploying additional fiber, which reduces the cost and time associated with civil works. This wavelength-based capacity model also permits services with differing rates and protocols to share the same optical infrastructure.

"Inherent topology flexibility with intelligent chromatic dispersion" is also an advantage because DWDM transport systems incorporate optical planning and dispersion-management capabilities that support wavelength routing across long-haul, metro, and meshed optical topologies. Chromatic-dispersion management preserves signal integrity as wavelengths traverse fiber spans, amplifiers, and optical nodes, enabling network designers to use available fiber paths efficiently while scaling capacity. Cisco describes DWDM as an optical transport technology that multiplexes multiple wavelengths onto a single fiber to increase bandwidth, while ITU-T G.872 defines the architecture and functionality of optical transport networks that support this wavelength-based approach. See [Cisco DWDM overview](#) and [ITU-T Recommendation G.872](#).

QUESTION NO: 22

As a part of a network design, you should tighten security to prevent man-in-the-middle. Which two security options ensure that authorized ARP responses take place according to known IP-to-MAC address mapping? (Choose two)

- A. DHCP snooping
- B. ARP spoofing
- C. ARP rate limiting
- D. Dynamic ARP Inspection

ANSWER: A D

Explanation:

DHCP snooping and Dynamic ARP Inspection work together to protect the LAN from ARP-based man-in-the-middle attacks. DHCP snooping classifies switch interfaces as trusted or untrusted and records DHCP lease information learned on untrusted access ports in a DHCP snooping binding database. Each binding associates a client IP address with its learned MAC address, VLAN, interface, and lease information. This creates the trusted IP-to-MAC mapping required for subsequent validation.

Dynamic ARP Inspection uses that DHCP snooping binding database to inspect ARP packets received on untrusted interfaces. For an ARP request or reply to be accepted, its sender address information must match the authorized binding for that VLAN and interface. Consequently, the switch can discard forged ARP messages that attempt to associate a legitimate IP address with an attacker-controlled MAC address. Trusted infrastructure-facing ports are configured appropriately so that legitimate DHCP and ARP control traffic can pass while access-port clients remain subject to validation.

In a Cisco campus design, enabling DHCP snooping first and then enabling Dynamic ARP Inspection on the relevant user VLANs provides the mapping source and the enforcement mechanism needed for this protection. Cisco documents the dependency and validation behavior in its [DHCP Snooping Configuration Guide](#) and [Dynamic ARP Inspection Configuration Guide](#).

QUESTION NO: 23

You are designing an optical network. Your goal is to ensure that your design contains the highest degree of resiliency. In which two ways should you leverage a wavelength-switched optical network solution in your network design? (Choose two.)

- A. a wavelength-switched optical network guarantees restoration based strictly on the shortest path available
- B. a wavelength-switched optical network provides fault tolerance for single failures only
- C. a wavelength-switched optical network takes linear and nonlinear optical impairment calculation into account
- D. a wavelength-switched optical network assigns routing and wavelength information
- E. a wavelength-switched optical network eliminates the need for dispersion compensating units in a network

ANSWER: C D

Explanation:

A wavelength-switched optical network takes linear and nonlinear optical impairment calculation into account because reliable optical service provisioning depends on validating that a proposed wavelength path has an acceptable optical signal-to-noise ratio, dispersion performance, power level, and nonlinear-impairment margin. Incorporating these constraints into path computation avoids provisioning a nominally available path that cannot sustain the required service quality, which directly improves availability and restoration success.

A wavelength-switched optical network assigns routing and wavelength information through routing and wavelength assignment. The control plane can determine an end-to-end optical route and select a compatible wavelength, while considering network topology, resource availability, and applicable optical constraints. This enables automated establishment and re-establishment of lightpaths after a failure, reducing restoration time and operational dependence on manual wavelength planning. The IETF WSON framework describes the control-plane information needed for wavelength availability, wavelength assignment, and optical-network path computation. Cisco optical transport platforms similarly use control-plane and planning functions to support resilient optical connectivity. See the [IETF WSON Framework](#) and Cisco's [NCS 2000 Configuration Guide](#).

QUESTION NO: 24

As part of network design, two geographically separated data centers must be interconnected using Ethernet-over-MPLS pseudowire. The link between the sites is stable, the topology has no apparent loops, and the root bridges for the respective VLANs are stable and unchanging. Which aspect must be the part of the design to mitigate the risk of connectivity issues between the data centers?

- A. Enable Spanning Tree on one data center, and Rapid Reconfiguration of Spanning tree on the other
- B. Ensure that the spanning tree diameter for one or more VLANs is not too large.
- C. Enable UDLD on the link between the data centers.
- D. Enable root guard on the link between the data centers.

ANSWER: B

Explanation:

Ensure that the spanning tree diameter for one or more VLANs is not too large. A Layer 2 Ethernet-over-MPLS pseudowire transparently extends a VLAN and its spanning-tree control plane between the data centers. Therefore, BPDUs must traverse every bridge hop in the resulting end-to-end Layer 2 topology. Even where physical links, topology, and root-bridge placement are stable, an excessive bridge diameter can make the spanning-tree timing budget inadequate. The root bridge's hello, maximum-age, and forward-delay timer values must allow BPDUs and topology information to propagate across the entire active tree before they age out. If the actual diameter exceeds the configured or assumed value, switches may age valid superior information prematurely, producing unnecessary reconvergence and temporary loss of forwarding connectivity. The design should calculate the diameter per VLAN—including the Layer 2 extension and all intermediate bridges—and use STP timer values appropriate for that diameter, while keeping the Layer 2 domain intentionally bounded. Cisco documents that the network diameter parameter represents the maximum switch-to-switch distance from the root and is used when calculating STP timer values. See [Cisco's Understanding Spanning Tree Protocol](#) and the [Cisco IOS XE STP configuration guide](#).

QUESTION NO: 25

Which option is a critical mechanism to optimize convergence speed when using MPLS

FRR?

- A. IGP timers
- B. Bandwidth reservation
- C. Shared risk link groups
- D. Down detection

ANSWER: D

Explanation:

Down detection is a critical mechanism for optimizing convergence with MPLS Fast Reroute because the repair process cannot begin until the router recognizes that the protected link, next hop, or node has failed. MPLS FRR provides a precomputed and pre-signaled backup path, allowing the point of local repair to switch traffic onto that backup immediately after a failure is detected. Consequently, rapid and reliable failure detection directly determines how quickly traffic is restored and is essential to achieving the commonly targeted sub-50-millisecond protection interval.

Fast physical-layer indications may detect some failures, but they do not cover every forwarding failure scenario. Bidirectional Forwarding Detection can provide much faster, protocol-independent detection of a failed forwarding path and is commonly used with MPLS TE and FRR deployments. Once the failure is declared, the local repair mechanism redirects labeled traffic to the established bypass or detour LSP without waiting for network-wide routing reconvergence. Cisco's overview of [Bidirectional Forwarding Detection](#) describes its role in rapid path-failure detection. The MPLS FRR behavior and local protection model are also standardized in [RFC 4090](#).

QUESTION NO: 26

Which three items do you recommend for control plane hardening of an infrastructure device? (Choose three)

- A. To enable unused services
- B. Warning banners
- C. Routing protocol authentication
- D. Control Plane Policing
- E. Redundant AAA servers
- F. SNMPv3

ANSWER: C D F

Explanation:

Control-plane hardening protects the router or switch CPU and the protocols used to administer or exchange critical state with the device. **Routing protocol authentication** is recommended because it validates routing peers and protects routing adjacencies from unauthorized or spoofed protocol updates. Cisco routing protocol security guidance identifies authentication as a fundamental protection for routing exchanges.

Control Plane Policing is essential because it classifies and rate-limits traffic destined for the device control plane. This helps preserve CPU resources and maintain device availability during excessive or malicious traffic conditions, while allowing explicitly defined control and management traffic to be handled appropriately.

SNMPv3 is the recommended SNMP version for secure device management because it supports user-based authentication and privacy (encryption). Using authenticated and encrypted management-plane access reduces exposure of monitoring and administrative information and helps ensure that SNMP requests come from authorized users. Together, these measures protect the routing-control function, constrain traffic delivered to the CPU, and secure a common management protocol. Cisco documentation on [routing protocol authentication](#) and [Control Plane Policing](#) provides implementation guidance.

QUESTION NO: 27

You are implementing a one-to-many multicast solution for a large service provider network. Which technology offers optimal routing of multicast traffic?

- A. PIM sparse mode
- B. PIM SSM
- C. Anycast RP
- D. MSDP
- E. Bidirectional PIM.

ANSWER: B

Explanation:

PIM SSM is correct because Source-Specific Multicast is designed specifically for one-to-many multicast delivery. Receivers explicitly join a channel identified by both a multicast group and a particular source, conventionally written as (S,G). This allows the network to build a shortest-path distribution tree directly from the known source to interested receivers. Consequently, multicast packets follow an efficient, source-rooted path without initial shared-tree forwarding or a rendezvous-point dependency.

For a large service-provider deployment, this model provides important operational scalability. Channel membership is source-specific, so the network does not need to maintain rendezvous-point state for the service or use interdomain source-discovery mechanisms. SSM also creates a clear receiver-to-source subscription model that fits common one-way services

such as IPTV, live video distribution, financial feeds, and software distribution. The IETF defines SSM as an approach in which receivers specify both the source and group, with distribution based on source-rooted shortest-path trees. Cisco's multicast configuration guidance likewise describes SSM as a PIM deployment model for source-specific channels and direct source-tree forwarding. See [RFC 4607: Source-Specific Multicast for IP](#) and [Cisco IOS XE SSM configuration guide](#).

QUESTION NO: 28

A large enterprise network has two data centers and a WLAN edge with a large hub-and-spoke network. The complete network is configured as a single OSPF area, and spoke routers are connected to unreliable WAN links. Which two changes should you make to deploy LSA on the spoke routers? (Choose two)

- A. Place spoke routers in stub areas
- B. Make the hub routers ABR
- C. Make the hub routers ASBR
- D. Place spoke routers in totally stubby areas
- E. Keep the spoke routers in normal areas

ANSWER: B D

Explanation:

Making the hub routers ABRs and placing spoke routers in totally stubby areas creates the required OSPF hierarchy for a hub-and-spoke deployment. The hub must function as an Area Border Router because it connects each nonbackbone spoke area to the backbone or other OSPF areas. This lets the hub control which routing information is advertised toward the spokes and provides the boundary at which area-specific LSA behavior is applied.

A totally stubby area is well suited to remote spokes connected across constrained or unreliable WAN links. It prevents external LSAs and interarea summary LSAs from being flooded into the spoke area, while the ABR injects a default route for destinations outside the area. Consequently, spoke routers maintain a substantially smaller link-state database and perform fewer SPF-related processing tasks, reducing routing-control traffic over the WAN while retaining reachability through the hub. Totally stubby areas are a Cisco-supported extension configured on the ABR and consistently on the routers within that area. Cisco documents OSPF area types and the ABR's role in route advertisement at [Cisco: OSPF Design Guide](#) and describes stub-area behavior in the [Cisco IOS OSPF Stub Areas documentation](#).

QUESTION NO: 29 - (DRAG DROP)

Drag and drop the NETCONF layers on the left onto their appropriate description on the right.

transport	defines a set of base protocol operations
messages	provides a communication path between the client and server
operations	provides a framing mechanism for encoding RPCs
content	holds information on data models and protocol operations

ANSWER:



Explanation:

NETCONF is organized as a layered protocol architecture so that the way a device is reached, the way requests are encoded, the standard actions available to a client, and the management information being exchanged remain clearly separated. The transport layer is responsible for the underlying client-server communication path. In practical deployments, this is commonly provided through a secure session such as Secure Shell, which gives NETCONF a reliable path over which management traffic can travel.

Once that path exists, the messages layer defines how NETCONF remote procedure call requests and replies are framed and encoded for transmission. This makes it possible for both endpoints to recognize the boundaries of each management message while sharing the transport session. The operations layer then defines the base protocol operations that a NETCONF client can invoke, such as retrieving or editing configuration and managing configuration datastores. Finally, the content layer is the actual management information carried by the protocol: the configuration and state data structured by data models, together with the operation-related content needed by the exchange.

This separation lets NETCONF use a consistent management protocol independently of the particular data model being managed or the secure transport carrying the session. Cisco NETCONF implementations follow this same model. See the [NETCONF Protocol RFC](#) for the protocol architecture and the [Cisco NETCONF documentation](#) for implementation context.

QUESTION NO: 30

Which mechanism provides fast path failure detection?

- A. Non-Stop Forwarding
- B. Carrier delay
- C. Graceful restart
- D. UDLD
- E. Fast hello packets
- F. iSPF

ANSWER: E

Explanation:

Fast hello packets provide fast path-failure detection by reducing the interval at which routing peers transmit and expect liveness messages. If consecutive hello packets are not received within the configured failure-detection interval, the neighbor relationship is declared down quickly. The routing protocol can then remove routes learned through that peer, run its convergence logic, and select an alternate path much sooner than it would with standard hello and dead timers.

This mechanism is commonly used where rapid routing convergence is important and the network can tolerate the additional control-plane traffic created by more frequent keepalives. For example, OSPF supports minimal dead intervals and a configurable hello multiplier, allowing several hello packets per second and thereby detecting loss of adjacency in

substantially less time than default timers. Fast hello operation detects loss of routing-peer reachability, which is the key signal needed to trigger reconvergence after a path failure.

Cisco documents OSPF fast hello behavior in its [OSPF Fast Hello Packets guide](#). The underlying OSPF neighbor liveness behavior is also specified in [RFC 2328, Section 10.5](#).

QUESTION NO: 31

Which two options are IoT use cases that require the low-latency and high reliability that

5G networks provide? (Choose two)

- A. Sports and Fitness
- B. Smart Home
- C. Automotive
- D. Smart Cities
- E. Industrial Automation
- F. Health and wellness

ANSWER: C E

Explanation:

Automotive and Industrial Automation are the IoT domains most directly associated with 5G ultra-reliable low-latency communications. Automotive systems can depend on rapid, predictable exchange of information for connected-vehicle safety functions, cooperative driving, and vehicle-to-everything communications. In these scenarios, delayed or lost data can affect safety-related decisions, so consistently low latency and highly dependable wireless connectivity are essential rather than merely beneficial.

Industrial Automation similarly uses communications to support time-sensitive control loops, coordinated robots, automated guided vehicles, machine safety, and real-time process control. A wireless network serving these functions must deliver data within strict timing bounds and maintain service reliability in environments where interruption can halt production or create operational hazards. 5G was designed to address this category of demanding communication through ultra-reliable low-latency capabilities, alongside support for high device density and broad IoT connectivity. The 3GPP overview describes 5G service capabilities and architecture at [3GPP 5G System Overview](#). Cisco also discusses the role of 5G in industrial use cases at [Cisco Industrial IoT](#).

QUESTION NO: 32 - (SIMULATION)

Drag the QoS tools on the left and drop each into its corresponding function on the right.

Policing	Addresses congestion that is due to speed mismatches when CIR is not exceeded.
Marking	Drops traffic to ensure that the committed or offered rate are not exceeded.
Buffering	Allows drops to be minimized based on traffic classification when CIR is exceeded.
WRED	Allows for consistent classification within a DiffServ domain.
Shaping	Avoids congestion via selective traffic dropping within the network.
ECN	Avoids congestion by end hosts reducing their traffic rates when congestion is detected.

ANSWER: Check the explanation for the answer

Explanation:

Match the QoS tools as follows:

Policing enforces a rate by dropping excess packets, while shaping smooths traffic for a slower downstream rate. WRED proactively and selectively drops packets to prevent queue congestion; ECN signals congestion without requiring packet drops.

QUESTION NO: 33

When you design a network that uses IPsec, where can you reduce MTU to avoid network fragmentation?

- A. on both ends of the TCP connection
- B. on the side closest to the client
- C. on the side closest to the server
- D. in the WAN

ANSWER: A

Explanation:

on both ends of the TCP connection is correct because IPsec encapsulation adds headers and, depending on the security protocol and tunnel mode, can add substantial overhead to each packet. A packet that was valid for the original path MTU can therefore exceed the effective MTU after it is protected by IPsec. Reducing the usable packet size at both TCP endpoints ensures that each endpoint sends segments that remain within the effective path MTU after IPsec overhead is added. This prevents avoidable fragmentation in either direction of a bidirectional TCP flow and avoids the performance and reliability issues associated with fragmented encrypted traffic.

In practice, this objective is often implemented by setting an appropriate IP MTU on tunnel-facing interfaces or by adjusting TCP maximum segment size handling so that advertised segment sizes account for encapsulation overhead. The selected value must accommodate the underlying transport MTU and the specific IPsec encapsulation used. Cisco describes how IPsec overhead and fragmentation behavior affect packet forwarding in its [IPsec fragmentation documentation](#). The underlying requirement to respect the path MTU is also defined in [RFC 1191](#).