

Cisco Certified Design Expert (CCDE v3.0)

Cisco 400-007

Version Demo

Total Demo Questions: 66

Total Premium Questions: 666

Buy Premium PDF

<https://passqueen.com>

support@passqueen.com

Topic Break Down

Topic	No. of Questions
Topic 1, Network Design	216
Topic 2, Network Infrastructure	159
Topic 3, Network Services	64
Topic 4, Security	112
Topic 5, Automation	31
Topic 6, Emerging Technologies	83
Topic 7, Mix Questions	1
Total	666

QUESTION NO: 1

Which two design solutions ensure sub-50 msec of the convergence time after a link failure in the network? (Choose two)

- A. BFD
- B. Ti-LFA
- C. Minimal BGP scan time
- D. MPLS-FRR
- E. IGP fast hello

ANSWER: B D

Explanation:

Ti-LFA and MPLS-FRR are designed to provide rapid, local traffic protection after a failure, avoiding the delay associated with waiting for complete control-plane reconvergence. Ti-LFA (Topology-Independent Loop-Free Alternate) is Segment Routing's fast-reroute mechanism. It precomputes a repair path that is loop-free for a protected link, adjacent node, or SRLG failure. When the failure is detected, the point of local repair immediately imposes the appropriate Segment Routing label stack and forwards traffic over the precomputed backup path. This local repair model supports sub-50-ms restoration when paired with suitably fast failure detection.

MPLS-FRR likewise provides local protection using pre-established backup LSPs. In facility backup mode, traffic can be rapidly redirected around a protected link or node through a bypass tunnel; in one-to-one backup mode, a detour protects an individual LSP. Because the backup forwarding state is already installed before the fault occurs, switchover does not depend on RSVP-TE signaling or IGP recomputation after the event. Cisco describes MPLS TE Fast Reroute as supporting recovery times of less than 50 milliseconds, and documents Ti-LFA as Segment Routing fast reroute with precomputed repair paths. See [Cisco Ti-LFA documentation](#) and [Cisco MPLS TE Fast Reroute documentation](#).

QUESTION NO: 2

Which solution component helps to achieve rapid migration to the cloud for SaaS and public cloud leveraging SD-WAN capabilities?

- A. service-oriented cloud architecture
- B. Cloud onramp
- C. cloud registry
- D. microservices in the cloud

ANSWER: B

Explanation:

Cloud onramp is the SD-WAN solution component that accelerates adoption of SaaS and public-cloud services. Cisco SD-WAN Cloud OnRamp extends SD-WAN policy and application-aware transport selection to cloud destinations, allowing enterprises to connect sites and users to cloud applications without designing separate, manually managed connectivity for every location. Cloud OnRamp for SaaS identifies supported SaaS applications and measures the performance of available paths, enabling traffic to be forwarded over the path that best meets application requirements. This preserves a consistent user experience as workloads and collaboration services move from private data centers to cloud-delivered applications.

For public-cloud workloads, Cisco's Cloud OnRamp capabilities simplify connectivity between the SD-WAN fabric and cloud environments while retaining centralized policy, segmentation, visibility, and automated provisioning. These functions reduce the operational time and complexity associated with migration because cloud access is incorporated into the existing SD-WAN architecture rather than treated as a separate network domain. The result is faster rollout of cloud connectivity with performance-aware routing and unified operations across branch, internet, and cloud paths. Cisco documents Cloud OnRamp as a cloud-connectivity capability of Cisco SD-WAN at [Cisco Cloud OnRamp](#) and provides configuration guidance at [Cisco SD-WAN Cloud OnRamp Configuration Guide](#).

QUESTION NO: 3

Company XYZ wants to improve the security design of their network to include protection from reconnaissance and DoS attacks on their sub interfaces destined toward next hop routers. Which technology can be used to prevent these types of attacks?

- A. MPP
- B. CPPr
- C. CoPP
- D. DPP

ANSWER: C

Explanation:

CoPP is the appropriate technology because it protects the router control plane from traffic directed at the router itself, including packets addressed to IP addresses configured on physical interfaces or subinterfaces. Reconnaissance activity, such as scans directed at router addresses, and denial-of-service floods can force packets to be handled by the route processor. Without protection, excessive control-plane traffic can consume CPU and memory resources, impair routing-protocol processing, and affect management-plane availability.

CoPP classifies traffic that is destined for the control plane and applies a service policy directly to that control plane. The policy can identify important protocol traffic, management traffic, and undesirable or excessive traffic, then police each class at an appropriate rate. This preserves route-processor capacity for essential functions while limiting the impact of malicious or unexpected packet rates. Cisco specifically describes CoPP as a mechanism for managing and protecting traffic destined to the router CPU, making it a core router-hardening control for this design requirement. See Cisco's [Control Plane Policing implementation guidance](#) and the [Cisco IOS XE CoPP configuration guide](#).

QUESTION NO: 4

An enterprise is implementing Resource Public Key Infrastructure validation on its Internet-facing BGP speakers. Which two statements about BGP origin validation are true? (Choose two.)

- A. A route is valid when a covering ROA authorizes both its origin AS and prefix length.
- B. A route is invalid when a covering ROA exists but authorizes a different origin AS.
- C. A route is invalid whenever no ROA exists for the advertised prefix.
- D. Origin validation verifies every autonomous system in the AS_PATH attribute.
- E. RPKI validation replaces the need for prefix filters from customers.

ANSWER: A B

Explanation:

A route is considered **valid** when a Route Origin Authorization covers the advertised prefix and the route origin AS matches the AS authorized by that ROA. The ROA must also authorize the advertised prefix length, including the configured maximum length where one is present. Validation gives the receiving BGP speaker cryptographically verifiable information about whether the origin AS is authorized to originate the prefix.

A route is considered **invalid** when a covering ROA exists but the route origin AS is not authorized, or when the advertised prefix length exceeds the maximum length permitted by the ROA. Invalid is distinct from not found: not found means that no covering ROA exists. Origin validation does not itself verify the complete AS path; it validates only the origin authorization. The RPKI-based origin-validation process is specified in [RFC 6811](#).

QUESTION NO: 5

Before migrating anything to the cloud, what are three cloud readiness assessment steps that are required to perform? (Choose three.)

- A. Determine the geolocation requirement
- B. Determine the technology obsolescence
- C. Identify the scope and business cases for migration.
- D. Assess infrastructure requirements.
- E. Evaluate available in-house resources
- F. Assess database security requirements
- G. Assess the cloud brokerage solutions

ANSWER: C D E

Explanation:

A cloud-readiness assessment begins by establishing why the organization is migrating and exactly which workloads, applications, and business outcomes are in scope. Therefore, *Identify the scope and business cases for migration*. is essential: it connects migration decisions to measurable business value, priorities, risk tolerance, and an appropriate migration approach for each workload. Cisco describes cloud migration as a transformation that should align technology decisions with business needs and operational outcomes.

Assess infrastructure requirements. is also required because workload placement depends on the current network, compute, storage, connectivity, performance, availability, and integration dependencies. This assessment identifies the technical foundation and any changes needed to support secure, reliable access to cloud-hosted services.

Finally, *Evaluate available in-house resources* establishes whether the organization has the skills, capacity, operating model, and governance needed to execute and sustain the migration. Readiness is not solely a platform decision; operational ownership and expertise determine whether cloud services can be deployed and managed effectively after migration. These assessment areas provide the business, technical, and operational baseline needed before committing workloads to a cloud environment. See Cisco's [cloud solutions overview](#) and its [cloud migration resources](#).

QUESTION NO: 6

Which BGP feature provides fast convergence?

- A. BGP PIC |
- B. BGP-EVPN
- C. BGP FlowSpec
- D. BGP-LS

ANSWER: A

Explanation:

BGP PIC (Prefix Independent Convergence) provides fast BGP convergence following a failure. Its central design objective is to minimize traffic loss by preparing backup forwarding information before a failure occurs. Rather than requiring the router to process and update every affected BGP prefix individually after a next-hop, link, or node failure, BGP PIC uses indirection in the forwarding plane: many prefixes can point to a common BGP next-hop or repair path. When that primary forwarding path fails, the router can rapidly switch the relevant forwarding entries to a precomputed alternate next hop. This substantially reduces FIB convergence time, particularly in large routing tables where conventional per-prefix updates could otherwise cause significant disruption. BGP PIC Core is commonly used to protect reachability through the provider core, while BGP PIC Edge supports rapid restoration for PE-facing services and external destinations. The feature is therefore

especially valuable in service-provider and large enterprise designs that have stringent availability requirements and need BGP-based forwarding to recover quickly from infrastructure failures. Cisco describes BGP PIC as a mechanism that installs backup paths in advance and enables rapid forwarding-plane updates upon failure. See the [Cisco IOS XE BGP PIC configuration guide](#) and Cisco's [BGP PIC overview and configuration documentation](#).

QUESTION NO: 7

Which two mechanisms avoid suboptimal routing in a network with dynamic mutual redistribution between multiple OSPFv2 and EIGRP boundaries? (Choose two.)

- A. AD manipulation
- B. Matching OSPF external routes
- C. Route tagging
- D. Route filtering
- E. Matching EIGRP process ID

ANSWER: A C

Explanation:

AD manipulation and **Route tagging** together provide control-plane safeguards where OSPFv2 and EIGRP are mutually redistributed at more than one location. Administrative distance establishes a deterministic preference when the same destination is available through routes learned by different protocols or through redistributed copies of a route. Deliberately setting distance values prevents a redistributed external path from being preferred over the intended native or preferred protocol path, which helps maintain consistent, optimal forwarding choices across redistribution boundaries.

Route tagging preserves route-origin information as a prefix crosses from one protocol into the other. A redistribution policy can set a tag when exporting a route and then match that tag on later redistribution operations. This allows the design to identify routes that have already crossed a boundary and enforce the intended redistribution policy, rather than allowing uncontrolled feedback between protocols. In a multi-boundary design, the combination of explicit route-origin metadata and predictable route preference is especially important because each boundary otherwise has only local knowledge of a route's history. Cisco documents route tagging as a redistribution-policy attribute and administrative distance as the route-selection mechanism used when multiple sources advertise a destination. See [Cisco IOS Administrative Distance documentation](#) and [Cisco IOS routing policy documentation](#).

QUESTION NO: 8

A Tier-3 Service Provider is evolving into a Tier-2 Service Provider due to the amount of Enterprise business it is receiving. The network engineers are re-evaluating their IP/MPLS design considerations in order to support duplicate/overlapping IP addressing from their Enterprise customers within each Layer3 VPN. Which concept would need to be reviewed to ensure stability in their network?

- A. Assigning unique Route Distinguishers
- B. Assigning unique Route Target IDs
- C. Assigning unique IP address space for the Enterprise NAT/Firewalls
- D. Assigning unique VRF IDs to each L3VPN

ANSWER: A

Explanation:

Assigning unique Route Distinguishers is correct because an MPLS Layer 3 VPN must be able to carry identical customer IPv4 prefixes without treating them as the same route in the provider's VPN control plane. A Route Distinguisher is an 8-byte value prepended to a customer IPv4 prefix to form a 96-bit VPN-IPv4 address. For example, two separate customer VPNs

may both advertise 10.0.0.0/8; unique Route Distinguishers make those advertisements distinct VPN-IPv4 NLRI and allow BGP to maintain and distribute them independently. This preserves route uniqueness in the provider network while allowing each customer to continue using its own private or overlapping addressing plan. Route Distinguishers identify routes; they do not determine VPN membership. VPN membership and route-sharing policy are applied separately through route-target extended communities. In a scalable provider design, Route Distinguisher allocation should therefore be planned so that VPN routes requiring independent identities receive distinct values, including where overlapping prefixes are present. RFC 4364 defines the VPN-IPv4 address as the combination of an RD and IPv4 prefix and describes its role in making overlapping address spaces unique. Cisco's MPLS Layer 3 VPN documentation likewise describes the RD as the mechanism that enables support for overlapping customer address spaces. See [RFC 4364](#) and [Cisco MPLS Layer 3 VPN overview](#).

QUESTION NO: 9

Company XYZ asks for design recommendations for Layer 2 redundancy. The company wants to prioritize fast convergence and resiliency elements. In the design. Which two technologies are recommended? (Choose two.)

- A. Design MLAG/MC-LAG into the network wherever possible.
- B. Configure DHCP snooping on the switches.
- C. Use root guard.
- D. Use BPDU guard.
- E. Use UniDirectional Link Detection.

ANSWER: A E

Explanation:

Design MLAG/MC-LAG into the network wherever possible and use UniDirectional Link Detection are the recommended technologies when the principal Layer 2 design goals are rapid convergence and resilience. Multi-chassis link aggregation allows a downstream device to form one logical aggregated connection to two separate upstream switches. Both links can forward actively, avoiding dependence on a blocked spanning-tree backup link during normal operation. If an individual member link or one upstream chassis fails, the remaining active member links can continue carrying traffic, reducing disruption and providing resilient forwarding behavior. Cisco implementations of this architecture include virtual PortChannel and related multichassis EtherChannel designs.

UniDirectional Link Detection adds protection against a failure mode that ordinary interface state detection may not immediately identify: a link that can receive traffic in one direction but cannot transmit correctly in the other. Detecting and disabling or otherwise withdrawing such a faulty link prevents persistent forwarding loops, blackholing, and unstable Layer 2 behavior. In an aggregated or redundant topology, this enables traffic to use healthy paths promptly and predictably. Together, multichassis aggregation supplies resilient active/active pathing, while UDLD protects that pathing from unidirectional physical-link faults. See Cisco's [UDLD Configuration Guide](#) and [Virtual PortChannel overview](#).

QUESTION NO: 10

An enterprise service provider is planning to migrate the customer network to MPLS to connect cloud applications. The customer network team and service provider team are analyzing all process (tows before live migration and implementation. Before planning the migration, what is a crucial task that must be executed?

- A. real-time process monitoring and maintenance
- B. impact forecasts and risk analysis
- C. application packaging and deployment
- D. impact analysis and application refactoring

ANSWER: B

Explanation:

impact forecasts and risk analysis is the crucial activity to complete before an MPLS migration is planned. A migration changes the WAN forwarding and service-delivery model, so the customer and provider must first establish how applications, traffic classes, site connectivity, routing dependencies, security controls, and operational processes will behave during and after the transition. Forecasting the impact includes determining expected bandwidth and latency requirements, identifying cloud-application performance sensitivities, defining availability objectives, and validating whether the target MPLS design has adequate capacity and resilience.

Risk analysis converts those findings into a controlled migration approach. It identifies possible failure scenarios, such as routing-policy errors, loss of reachability to cloud services, overlapping addressing, QoS misclassification, or incomplete site cutovers. The teams can then define mitigations, maintenance windows, pilot sites, acceptance tests, rollback procedures, and clear responsibilities before any production changes occur. This assessment-led approach supports a design that meets business and technical requirements while reducing the likelihood and blast radius of a live-migration incident. Cisco's MPLS design guidance emphasizes planning services, traffic engineering, resiliency, and operational requirements as part of an MPLS deployment design; see [Cisco MPLS solutions](#) and [Cisco IOS XE MPLS Layer 3 VPN configuration guide](#).

QUESTION NO: 11

: 487

Which layer of the SDN architecture orchestrates how the applications are given the resources available in the network?

- A. orchestration layer
- B. northbound API
- C. control layer
- D. southbound API

ANSWER: C

Explanation:

The **control layer** is correct because it provides the centralized decision-making function in a software-defined networking architecture. The SDN controller, which operates in this layer, maintains an abstracted, network-wide view of topology, device state, forwarding information, and available capabilities. It uses that view to translate an application's intent or policy into the network behavior and resource allocations required to fulfill it. For example, when an application requests connectivity with specified bandwidth, latency, reachability, or security characteristics, the control layer determines suitable paths and programs the relevant forwarding behavior across the infrastructure. This separation lets applications express what they require without needing to calculate device-specific configurations or directly manage individual switches and routers. The controller also applies policy consistently and adapts allocations as topology, demand, or operational state changes. Cisco describes SDN as separating network control from forwarding and centralizing control logic in software, enabling policy-driven automation and programmability. See Cisco's [Software-Defined Networking overview](#) and the [Cisco explanation of SDN](#).

QUESTION NO: 12

Which two impacts of adding the IP event dampening feature to a network design are true? (Choose two.)

- A. It protects against routing loops.
- B. It switches traffic immediately after a link failure.
- C. It speeds up link failure detection.
- D. It reduces the utilization of system processing resources.
- E. It improves overall network stability.

ANSWER: D E

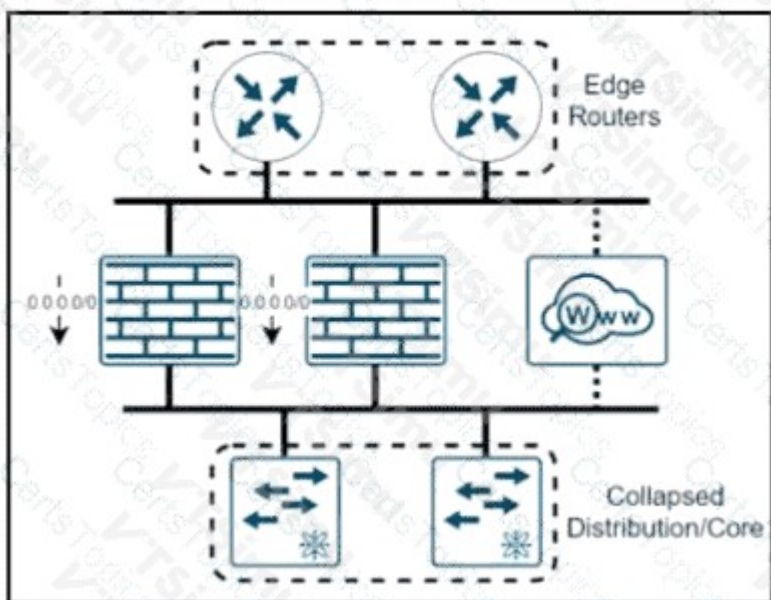
Explanation:

IP Event Dampening provides a penalty-based mechanism for managing unstable, repeatedly changing IP events, such as interface state transitions. When an event flaps often enough to exceed the configured suppress threshold, the device suppresses further processing of that event until its penalty decays below the reuse threshold. As a result, *It reduces the utilization of system processing resources*. Repeated link-state changes can otherwise initiate recurring control-plane processing, including interface-event handling, routing-protocol reactions, logging, and network-management notifications. Dampening limits this repetitive work while the affected condition remains unstable.

It improves overall network stability. By preventing a flapping event from repeatedly propagating through the device and triggering successive control-plane reactions, the feature reduces churn and makes network behavior more predictable during intermittent failures. This is a deliberate stability trade-off: dampening suppresses instability rather than accelerating fault detection or immediate recovery. Cisco describes IP Event Dampening as a method to suppress flapping events based on configurable penalty, suppress, reuse, and half-life values. See the [Cisco IOS XE IP Event Dampening Configuration Guide](#) and Cisco's [Route Flap Dampening Overview](#) for the underlying dampening concepts.

QUESTION NO: 13

Refer to the exhibit.



A company named XYZ needs to apply security policies for end-user browsing by installing a secure web proxy appliance. All the web traffic must be inspected by the appliance, and the remaining traffic must be inspected by an NGFW that has been upgraded with intrusion prevention system functionality. In which two ways must the routing be performed? (Choose two)

- A. Policy-based routing on the collapsed core
- B. Policy-based routing on the internet edge
- C. Policy-based routing on firewalls
- D. Static routing on the appliance

ANSWER: A D

Explanation:

Policy-based routing on the collapsed core is required because ordinary destination-based routing cannot distinguish web flows from other Internet-bound flows when they share the same destination prefixes or default route. The collapsed core is the appropriate insertion point, before traffic reaches the security devices. A policy can match the relevant web traffic,

such as HTTP and HTTPS flows, and set the secure web proxy as the next hop. This deterministically ensures that browsing traffic is sent first to the appliance that applies the organization's web-use and content-security policy, while traffic not matching the web policy follows the normal path toward NGFW and IPS inspection. Cisco documents that policy-based routing can classify packets using access-control-list criteria and override the normal routing-table next-hop decision: [Cisco IOS XE Policy-Based Routing Configuration Guide](#).

Static routing on the appliance is also needed to provide a defined forwarding and return path after proxy inspection. The proxy needs explicit reachability toward the NGFW or its designated transit interface so that inspected sessions continue through the firewall and IPS before Internet egress. A static route makes that service chain predictable and helps retain symmetric flow handling through the security infrastructure. Cisco describes static routes as manually configured routes used to specify a next hop or outgoing interface: [Cisco IOS XE Static Routes Configuration Guide](#).

QUESTION NO: 14

Which two actions must merchants do to be compliant with the Payment Card Industry Data Security Standard (PCI DSS)? (Choose two.)

- A. Conduct risk analyses
- B. Install firewalls
- C. Use antivirus software
- D. Establish monitoring policies
- E. Establish risk management policies

ANSWER: B C

Explanation:

Install firewalls is correct because PCI DSS requires organizations to define and implement network security controls that protect the cardholder data environment. These controls must restrict inbound and outbound network traffic to what is necessary, protect system components from untrusted networks, and be configured, maintained, and reviewed. Firewalls are a common implementation of these required network security controls and remain fundamental to PCI DSS compliance.

Use antivirus software is also correct because PCI DSS requires protection of systems against malicious software. Organizations must deploy anti-malware mechanisms where they are applicable, keep them current, perform periodic scans or continuous monitoring, and ensure the mechanisms cannot be disabled or altered improperly. Antivirus software is an established anti-malware control that supports these requirements, particularly for systems commonly targeted by malware.

These actions map directly to core PCI DSS technical requirements for protecting account data: network security controls under Requirement 1 and protection against malicious software under Requirement 5. The PCI Security Standards Council's current requirements and supporting guidance are available in the [PCI SSC Document Library](#) and the [PCI SSC merchant resources](#).

QUESTION NO: 15

The General Bank of Greece plans to upgrade its legacy end-of-life WAN network with a new flexible, manageable, and scalable solution. The main requirements are ZTP support, end-to-end encryption, application awareness, and segmentation. The CTO states that the main goal of the bank is CAPEX reduction. Which WAN technology should be used for the solution?

- A. SD-branch
- B. DMVPN with PfR
- C. Managed SD-WAN
- D. SD-WAN

ANSWER: D

Explanation:

SD-WAN is the appropriate WAN technology because it natively combines centralized policy and management with a secure overlay that can use diverse underlay transports, including lower-cost Internet broadband alongside private circuits where needed. This enables the bank to reduce capital expenditure by avoiding a design based solely on expensive dedicated WAN links, while retaining the ability to select paths according to application and business intent. Cisco SD-WAN provides Zero-Touch Provisioning, allowing branch devices to securely join the fabric and receive their configuration from centralized controllers rather than requiring extensive onsite staging. Cisco documents the ZTP workflow at [Cisco SD-WAN Zero-Touch Provisioning](#).

The SD-WAN fabric also establishes encrypted IPsec tunnels between WAN edge devices, protecting traffic across untrusted transports. Its application-aware routing capabilities classify applications and steer traffic using policy and measured path characteristics, rather than relying only on destination prefixes. In addition, VPN-based segmentation separates business domains and supports distinct policies across the same physical infrastructure. These capabilities are described in Cisco's [SD-WAN Security Overview](#). Thus, SD-WAN directly satisfies the operational, security, application, segmentation, scalability, and cost objectives stated for the bank.

QUESTION NO: 16

The network designer needs to use GLOP IP addresses to make them unique within their ASN. Which multicast address range will be considered?

- A. 239.0.0.0 to 239.255.255.255
- B. 224.0.0.0 to 224.0.0.255
- C. 233.0.0.0 to 233.255.255.255
- D. 232.0.0.0 to 232.255.255.255

ANSWER: C

Explanation:

The correct multicast range is **233.0.0.0 to 233.255.255.255**. GLOP addressing, specified by RFC 2770, uses the 233/8 IPv4 multicast block to derive globally unique multicast group addresses from an autonomous system number. Under the original GLOP scheme, a 16-bit ASN is encoded into the second and third octets of an address in 233/8. The remaining fourth octet provides 256 multicast group addresses for that autonomous system. For example, an ASN represented by the two octets 1 and 44 would use the block 233.1.44.0/24. This deterministic mapping enables an organization to select multicast groups associated with its ASN without requiring a separate, per-group allocation or coordination process. The range is therefore the applicable choice when the design requirement is GLOP-derived multicast addressing. RFC 2770 defines both the 233/8 allocation and the ASN-to-address mapping method: [RFC 2770](#). The IANA IPv4 multicast-address registry also identifies 233.0.0.0/8 as the GLOP address range: [IANA Multicast Addresses](#).

QUESTION NO: 17

Which two foundational aspects of IoT are still evolving and being worked on by the industry at large? (Choose two)

- A. WiFi protocols
- B. Regulatory domains
- C. Low energy Bluetooth sensors
- D. IoT consortia
- E. Standards

ANSWER: D E

Explanation:

IoT consortia and **Standards** are correct because broad IoT adoption depends on industry collaboration to define interoperable architectures, data models, connectivity approaches, device lifecycle practices, and security requirements. IoT spans many vendors and operational environments, so no single supplier can establish all of the conventions needed for devices, platforms, applications, and networks to work together consistently. Industry consortia provide venues in which vendors, operators, and technology organizations develop common specifications, reference architectures, certification programs, and implementation guidance. For example, the Open Connectivity Foundation publishes interoperability specifications and certification resources intended to enable devices from different vendors to communicate reliably.

Standards likewise remain an active area of development because IoT technology continually introduces new deployment models, protocols, threat considerations, and regulatory expectations. Mature IoT designs therefore track standards for secure device identification, software updates, vulnerability handling, data protection, and interoperable communications rather than treating standardization as complete. NIST's IoT cybersecurity guidance illustrates this continuing industry need by defining baseline capabilities that organizations should consider when procuring and deploying IoT devices. These evolving standards and collaborative consortia are foundational to scalable, multivendor, and secure IoT ecosystems.

References: [Open Connectivity Foundation Specifications](#); [NISTIR 8259: Foundational Cybersecurity Activities for IoT Device Manufacturers](#).

QUESTION NO: 18

Which two factors provide multifactor authentication for secure access to applications and data? (Choose two.)

- A. Persona-based
- B. Power-based
- C. Push-based
- D. Possession-based
- E. Pull-based
- F. Knowledge-based

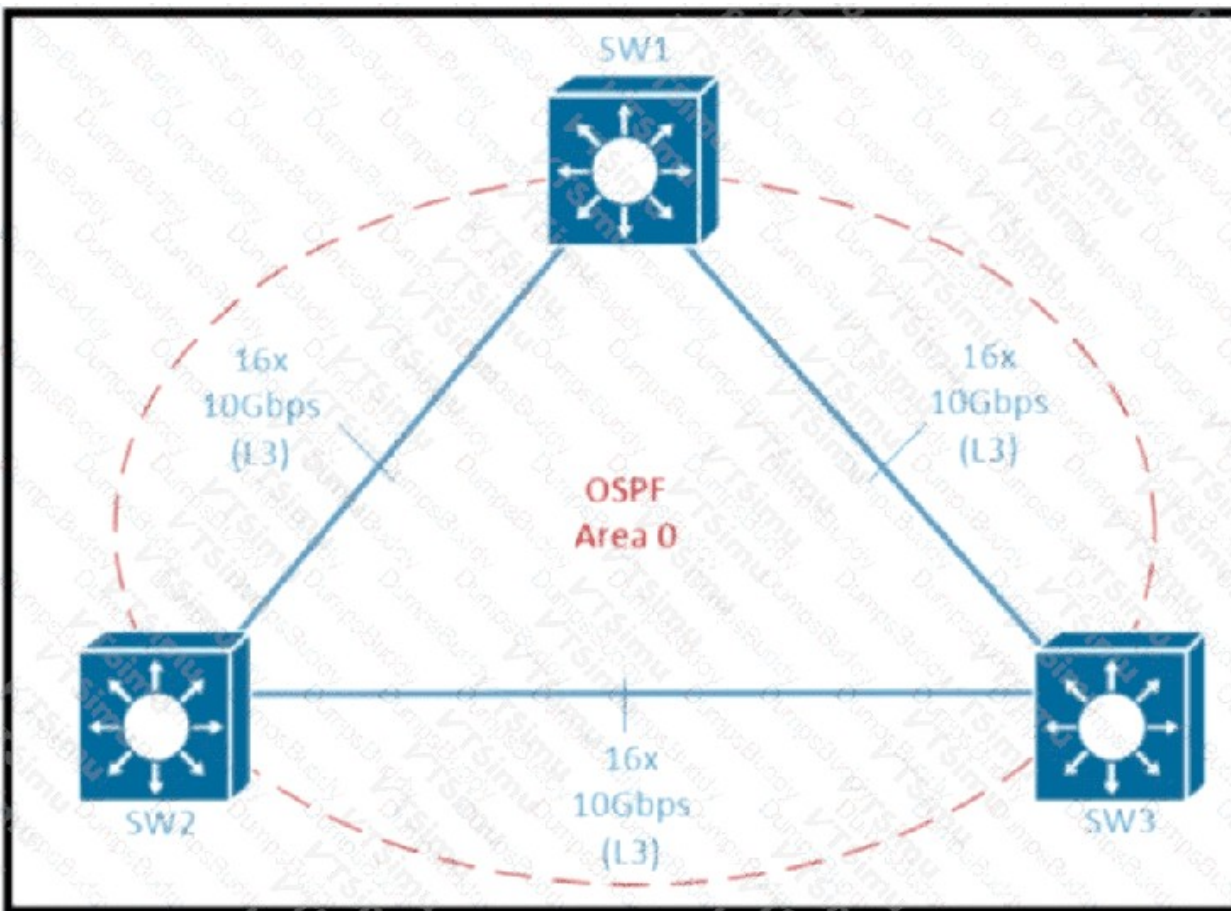
ANSWER: D F

Explanation:

Possession-based and **Knowledge-based** are distinct authentication-factor categories that can be combined to implement multifactor authentication. A possession factor verifies control of something assigned to or held by the user, such as a hardware token, smart card, security key, or registered authenticator device. A knowledge factor verifies information that the user is expected to know, such as a password, PIN, or passphrase. Requiring one of each creates stronger assurance than relying on either category alone, because successful authentication requires both the credential or secret and access to the associated device or token.

NIST's Digital Identity Guidelines define memorized secrets as knowledge-based authenticators and distinguish them from possession-based authenticators. This separation of authenticator types is fundamental to MFA design: the authenticators used must come from different factor categories. Cisco's MFA guidance likewise describes MFA as using two or more independent verification methods, commonly including something a user knows and something the user has. See [NIST SP 800-63B](#) and [Cisco: What Is Multi-Factor Authentication?](#).

QUESTION NO: 19



Refer to the exhibit Which two design options reduce the size of OSPF database in the shown topology? (Choose two.)

- A. Loop Free Alternate
- B. type 3 LSA filtering
- C. prefix suppression
- D. Layer 2 link aggregation between core switches
- E. incremental SPF

ANSWER: B C

Explanation:

type 3 LSA filtering reduces the link-state database held by routers in a receiving OSPF area by preventing selected interarea summary prefixes from being advertised into that area. Type 3 summary LSAs are generated by an Area Border Router to describe destinations from other areas. Applying filtering at the appropriate area boundary limits the number of those summaries that downstream routers must store and process. This is particularly useful when an area does not need reachability information for every prefix originated elsewhere, while a default route or selected aggregate routes provide the required connectivity. Cisco documents the role of interarea summary LSAs and ABR behavior in its [OSPF design documentation](#).

prefix suppression reduces OSPF LSDB scale by suppressing unnecessary transit-link prefixes from router LSAs. In a routed core, routers need the topology information needed for SPF, but they may not need every point-to-point or transit subnet to be advertised as a reachable destination prefix. Prefix suppression retains the necessary link-state topology while removing address-prefix information that is not needed for forwarding, reducing both LSDB content and route-table entries. Cisco describes this feature and its OSPF advertisement behavior in the [OSPF Prefix Suppression configuration guide](#).

QUESTION NO: 20

Which two descriptions of CWDM are true? (Choose two)

- A. Typically used over long distances, but requires optical amplification
- B. Uses the 850nm band
- C. Allows up to 32 optical carriers to be multiplexed onto a single fiber
- D. Shares the same transmission window as DWDM
- E. Passive CWDM devices require no electrical power

ANSWER: D E

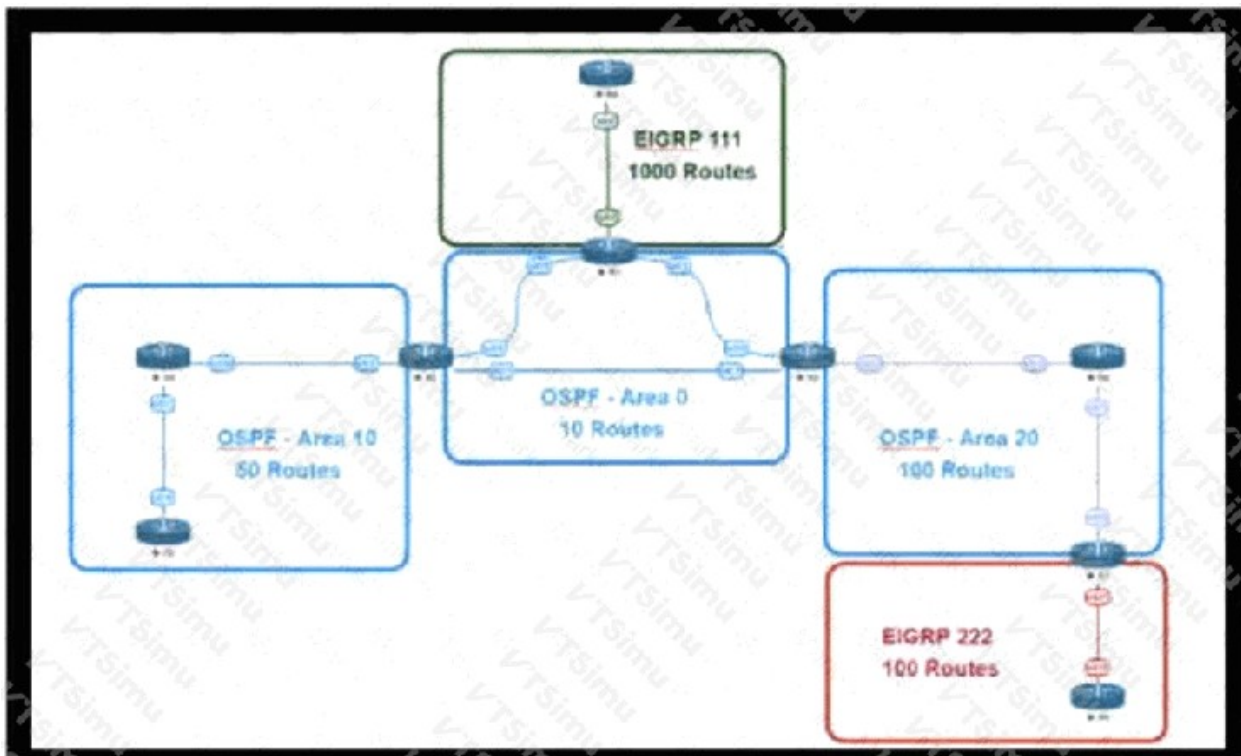
Explanation:

“Shares the same transmission window as DWDM” is correct in the practical sense that CWDM and DWDM are both wavelength-division multiplexing technologies operating over the low-loss optical-fiber wavelength regions used for optical transport. CWDM follows the ITU-T coarse-wavelength grid, whose nominal wavelengths span 1271 nm through 1611 nm with 20-nm channel spacing. This range includes the principal optical bands in which dense wavelength systems commonly operate, although DWDM normally concentrates its much more closely spaced channels in narrower portions of that overall fiber transmission region. Cisco CWDM optics are specified at the standardized CWDM wavelengths and are designed to carry multiple optical channels over a single single-mode fiber.

“Passive CWDM devices require no electrical power” is also correct. CWDM multiplexers and demultiplexers are passive optical components: they combine several wavelength-specific optical signals onto one fiber or separate those wavelengths at the far end. Because that wavelength-selection function is optical rather than electronic, the passive mux/demux unit itself does not need a power supply. This supports relatively simple, cost-conscious deployments, while the connected transceivers and any separately deployed active optical equipment can still require power. Cisco describes CWDM as a means of increasing fiber capacity through wavelength multiplexing, and the standardized coarse grid is defined by ITU-T. See [Cisco CWDM SFP data sheet](#) and [ITU-T G.694.2](#).

QUESTION NO: 21

Refer to the exhibit.



This network is running OSPF and EIGRP as the routing protocols. Mutual redistribution of the routing protocols has been configured on the appropriate ASBRs. The OSPF network must be designed so that flapping routes in EIGRP domains do not affect the SPF runs within OSPF. The design solution must not affect the way EIGRP routes are propagated into the EIGRP domains. Which technique accomplishes the requirement?

- A. route summarization at the ASBR interfaces facing the OSPF domain
- B. route summarization on the appropriate ASBRs
- C. route summarization on the appropriate ABRs
- D. route summarization on EIGRP routers connecting toward the ASBR

ANSWER: A

Explanation:

route summarization at the ASBR interfaces facing the OSPF domain is correct because the redistribution boundary is the appropriate place to hide EIGRP route instability from the OSPF domain. The ASBR advertises a stable aggregate for the redistributed EIGRP prefixes into OSPF rather than advertising each individual EIGRP subnet. As long as at least one component route remains available, a flap of an individual EIGRP route does not require the aggregate to be withdrawn or reoriginated in OSPF. This substantially reduces external LSA churn and prevents individual EIGRP reachability events from driving OSPF topology-processing activity throughout the OSPF domain.

The summarization is applied specifically toward OSPF at the redistribution point, so it changes only the information exported into OSPF. EIGRP continues to advertise and propagate its original, more-specific routes within its own domains; therefore, its internal routing behavior and visibility are preserved. Cisco documents ASBR external-route aggregation with the OSPF `summary-address` mechanism, which advertises an aggregate in place of matching redistributed external routes. OSPF's AS-external LSA behavior and ASBR role are defined in [RFC 2328](#); Cisco configuration guidance for OSPF route summarization is available in the [Cisco IOS OSPF Route Summarization documentation](#).

QUESTION NO: 22

A company uses equipment from multiple vendors in a data center fabric to deliver SDN, enable maximum flexibility, and provide the best return on investment. Which YANG data model should be adopted for comprehensive features to simplify and streamline automation for the SDN fabric?

- A. proprietary
- B. OpenConfig
- C. native
- D. IETF

ANSWER: B

Explanation:

OpenConfig is the appropriate choice for a heterogeneous SDN data center fabric because it provides vendor-neutral YANG models intended to create consistent configuration and operational-state interfaces across network platforms. This common modeling approach lets automation systems use the same data structures and semantics for comparable routing, interface, policy, and telemetry functions, rather than requiring distinct workflows for every device operating system. That consistency reduces integration effort, improves portability of automation code, and supports the flexibility and return-on-investment objectives stated in the scenario.

OpenConfig models are also operationally focused: they define both configuration and state data, supporting modern telemetry-driven automation and continuous validation workflows. For a fabric environment, this enables controllers and orchestration tools to retrieve normalized state, assess intended versus actual behavior, and automate changes with less vendor-specific abstraction. The OpenConfig project publishes the models openly and maintains them through industry collaboration, making them well suited to multi-vendor deployments that need broad, practical feature coverage. See the [OpenConfig project](#) and its [public YANG model repository](#) for the available models and their implementation-oriented design.

QUESTION NO: 23

A service provider is considering BGP Graceful Restart for route reflectors that serve a large VPN deployment. Which two statements about BGP Graceful Restart are true? (Choose two.)

- A. It can preserve forwarding while a BGP speaker restarts its control plane.
- B. A capable peer can retain routes from the restarting speaker as stale during the restart interval.
- C. It guarantees that no packet loss occurs during every hardware failure.
- D. It removes the requirement for redundant route reflectors.
- E. It requires BGP routes to be redistributed into the IGP before a restart.

ANSWER: A B

Explanation:

BGP Graceful Restart allows a restarting BGP speaker to preserve forwarding state while its control plane restarts, reducing traffic disruption when the relevant forwarding entries remain valid. A peer that supports the capability can retain routes received from the restarting speaker as stale during the negotiated restart period rather than immediately withdrawing them when the BGP session resets.

The restarting speaker and its peers advertise Graceful Restart capability information during BGP capability negotiation. The peer uses this information, together with the negotiated restart time and address-family forwarding-state indication, to determine whether stale routes can be retained. Graceful Restart is not a replacement for redundant route reflectors or resilient forwarding design; it is intended to reduce disruption during a planned or recoverable control-plane restart. The protocol behavior is defined in [RFC 4724](#).

QUESTION NO: 24 - (DRAG DROP)

The network team in XYZ Corp wants to modernize their infrastructure and is evaluating an implementation and migration plan to allow integration MPLS-based, Layer 2 Ethernet services managed by a service provider to connect branches and remote offices. To decrease OpEx and improve

response times when network components fail, XYZ Corp decided to acquire and deploy new routers. The network currently is operated over E1 leased lines (2 Mbps) with a managed CE service provided by the telco.

Drag and drop the implementation steps from the left onto the corresponding targets on the right in the correct order.

Activation of the MPLS-based L2VPN service.	Step 1
Monitor performance of the Ethernet service for a predefined period of time.	Step 2
Decommission E1 leased line.	Step 3
Deployment of new branch routers.	Step 4
Test and validation of SLA parameters in the L2VPN.	Step 5
Migrate E1 to the new Ethernet service.	Step 6

ANSWER:

Explanation:

A sound migration plan preserves service continuity by preparing the customer site first, validating the replacement service before depending on it, and retaining the old circuit until the new environment has demonstrated stable production behavior. Therefore, the sequence starts with **Deployment of new branch routers**. These routers replace the managed CE capability and must be physically installed, configured, reachable, and ready to connect to the provider handoff before the branch can use the target transport service.

Next comes **Activation of the MPLS-based L2VPN service**, which makes the provider-delivered Ethernet connectivity available to the newly deployed CE routers. Once the service is active, **Test and validation of SLA parameters in the L2VPN** confirms that the actual service characteristics align with the intended service design. Validation should cover the operational performance values relevant to the service, such as availability, delay, loss, and throughput, according to the provider agreement. Cisco's guidance on service-level validation and measurement is consistent with verifying service behavior before moving business traffic: [Cisco Service Level Agreement overview](#).

After the new service has passed validation, **Migrate E1 to the new Ethernet service** moves live branch connectivity from the existing 2-Mbps leased line to the MPLS Layer 2 VPN. The Ethernet service is then observed through **Monitor performance of the Ethernet service for a predefined period of time**. This stabilization period gives the team evidence that the new path is reliable under real traffic conditions and allows time to identify operational issues while the legacy circuit remains available as a safeguard.

Finally, **Decommission E1 leased line** removes the legacy service only after the migration and monitoring stages are complete. This ordering supports a controlled cutover, avoids an unnecessary outage risk, and enables the intended reduction in operational effort after the new branch infrastructure and provider service are proven in production.

QUESTION NO: 25

Which two impacts of adding the IP event dampening feature to a network design are true? (Choose two.)

- A. It protects against routing loops.
- B. It switches traffic immediately after a link failure.
- C. It speeds up link failure detection.
- D. It reduces the utilization of system processing resources.
- E. It improves overall network stability.

ANSWER: D E

Explanation:

It reduces the utilization of system processing resources and it improves overall network stability. IP Event Dampening is a Cisco IOS feature that monitors interface state transitions and, after an interface exceeds a configured flap threshold, suppresses routing-protocol notifications for that interface for a specified period. This prevents repeated up/down events on an unstable link from continually triggering route withdrawals, advertisements, SPF or other routing recalculations, and associated control-plane activity. Consequently, routers spend fewer CPU and memory resources processing churn caused by a flapping interface. Suppression also makes routing behavior more predictable because the network is insulated from transient link-state changes until the interface remains stable long enough for its penalty to decay below the reuse threshold. This is particularly valuable where a single unstable circuit could otherwise propagate frequent topology changes across a large routed domain. The design trade-off is intentionally delayed restoration of routing information for a suppressed interface, so dampening parameters should be selected carefully to balance stability and convergence requirements. Cisco documents IP Event Dampening as a mechanism for suppressing route updates resulting from interface flapping and reducing the effects of route flaps on the network. See [Cisco IOS IP Application Services Command Reference](#) and [Cisco IP Event Dampening Configuration Guide](#).

QUESTION NO: 26

multinational corporation with offices in various regions worldwide are looking for a network architecture that provides a balance between low- latency and cost-effectiveness for connecting their offices to the cloud. Which cloud connectivity option is the best fit for this corporation?

- A. DMVPN
- B. SDWAN SDCI
- C. MPLS direct connect
- D. cloud on-ramp

ANSWER: D

Explanation:

cloud on-ramp is the best fit because it is designed to provide optimized, scalable connectivity from distributed enterprise sites to cloud services. In a Cisco SD-WAN architecture, Cloud OnRamp extends the WAN fabric into cloud environments and can steer traffic over the most suitable available path based on application policy, performance measurements, and cloud-provider integration. This helps offices in different global regions reach cloud workloads or SaaS services with lower latency than a generic Internet path while retaining the flexibility to use cost-effective local Internet transport.

The approach is particularly appropriate for a multinational organization because it avoids requiring every site to rely solely on expensive private WAN circuits or on manually engineered cloud connections. Centralized SD-WAN policy can apply consistent segmentation, security, path selection, and application-aware routing across regions. Cloud OnRamp can also use cloud gateways and provider-adjacent connectivity to improve the path between the WAN and cloud services, giving the organization a practical balance of user experience, operational simplicity, and transport cost. Cisco describes Cloud OnRamp capabilities as part of its SD-WAN cloud-connectivity architecture in its [Cloud OnRamp for SaaS overview](#) and [Cisco SD-WAN documentation and solution resources](#).

QUESTION NO: 27

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

- A. LISP
- B. OTV
- C. VPLS
- D. EVPN

ANSWER: C

Explanation:

VPLS is correct because it emulates a multipoint Ethernet LAN across an MPLS provider backbone and uses traditional bridge-style MAC-address learning. When a VPLS provider-edge device receives an Ethernet frame from an attached customer interface, it records the source MAC address and the ingress attachment circuit in its forwarding database. If the destination MAC address is unknown, or the frame is broadcast or applicable multicast traffic, the device replicates and floods the frame to the other VPLS pseudowires. When return traffic arrives, source-MAC learning populates the forwarding information needed for later known-unicast forwarding. This behavior is commonly called flood-and-learn because data-plane traffic drives MAC reachability discovery, just as it does on a conventional Ethernet switch.

For DCI design, this characteristic matters because the amount of unknown-unicast, broadcast, and multicast replication can grow with the number of VPLS endpoints. Designers must therefore consider MAC-table scale, split-horizon behavior between pseudowires, and mechanisms that limit unnecessary flooding. Cisco’s VPLS documentation describes VPLS as an Ethernet multipoint service and identifies MAC-address learning and flooding as fundamental forwarding functions. The IETF VPLS architecture likewise specifies flooding and source-MAC learning as core elements of VPLS operation. See [Cisco VPLS Configuration Guide](#) and [RFC 4762: Virtual Private LAN Service Using LDP Signaling](#).

QUESTION NO: 28

In a distributed cloud-native environment, calls to services and cloud resources can fail caused by unanticipated events that will require longer periods of time to resolve. These faults can range in severity from a partial loss of connectivity to the complete failure of a service. In these situations, it ' s pointless for an application to continually retry an operation that is unlikely to succeed. Which pattern can prevent an application from repeatedly trying to execute an operation that ' s likely to fail?

- A. circuit breaker
- B. bulkhead
- C. fallback
- D. timeout

ANSWER: A

Explanation:

The **circuit breaker** pattern prevents an application from continually invoking an operation when a dependent service is experiencing persistent failures. The application monitors failures from calls to a remote service or cloud resource. Once failures exceed a defined threshold, the circuit breaker moves to an open state and immediately rejects subsequent calls rather than allowing requests to wait, time out, or consume resources on operations that are unlikely to succeed. This behavior protects the calling application, avoids amplifying an outage through excessive retries, and gives the failing dependency time to recover. After a configured recovery interval, the circuit breaker can enter a half-open state and permit a limited number of test requests. If those requests succeed, normal traffic can resume; if they fail, the circuit returns to the open state. This pattern is particularly important in distributed systems because failures can be transient, partial, or prolonged, and cascading resource exhaustion can otherwise spread across services. Cisco cloud-native designs commonly use this resilience approach through service-mesh or application-level traffic-management capabilities. See the [Microsoft Azure Architecture Center circuit breaker pattern](#) and the [Martin Fowler Circuit Breaker](#) reference.

QUESTION NO: 29

A route reflector serves clients located in multiple geographic regions. The clients receive a best path selected according to the route reflector's IGP cost to the BGP next hop, resulting in suboptimal traffic exit points for some regions. Which two statements describe Optimal Route Reflection? (Choose two.)

- A. It allows a route reflector to select a best path using an IGP location associated with a client or client group.
- B. It prevents route reflectors from advertising BGP paths to iBGP clients.
- C. It can reduce suboptimal routing caused by a route reflector selecting paths from its own topological location.
- D. It replaces the need for IGP reachability to BGP next hops.
- E. It requires every client to become a route reflector.

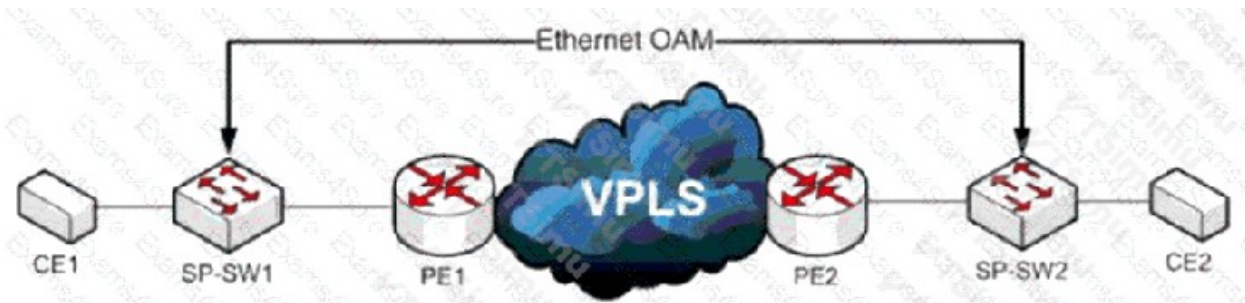
ANSWER: A C

Explanation:

Optimal Route Reflection enables a route reflector to select and advertise a best path from the perspective of an individual client or client group rather than exclusively from the route reflector's own location. This helps prevent a centralized route reflector from causing traffic to exit through an egress that is optimal for the reflector but distant from the receiving client.

ORR uses an IGP location, also called a primary or secondary path-selection location, to calculate the client-specific distance to BGP next hops. It does not alter the BGP path-selection rules themselves; it changes the IGP topology perspective used for the next-hop-cost comparison. Cisco documents this behavior in its [Optimal Route Reflection overview](#).

QUESTION NO: 30



Refer to the exhibit. A service provider has a requirement to use Ethernet OAM to detect end-to-end connectivity failures between SP-SW1 and SP-SW2. Which two ways to design this solution are true? (Choose two)

- A. Enable unicast heartbeat messages to be periodically exchanged between MEPs
- B. Enable Connectivity Fault Management on the SP switches
- C. Use upward maintenance endpoints on the SP switches
- D. Forward E-LMI PDUs over VPLS
- E. Forward LLDP PDUs over the VPLS

ANSWER: B C

Explanation:

Ethernet Connectivity Fault Management (CFM), specified by IEEE 802.1ag and incorporated into ITU-T Y.1731 Ethernet OAM functions, is appropriate for proactive end-to-end Layer 2 continuity monitoring. CFM defines a maintenance association spanning the service and places Maintenance End Points (MEPs) at its boundaries. The MEPs periodically transmit Continuity Check Messages (CCMs); loss of expected CCMs declares a connectivity fault. Therefore, enabling Connectivity Fault Management on the SP switches provides the OAM framework needed to monitor the Ethernet service between SP-SW1 and SP-SW2.

Using upward maintenance endpoints on the SP switches is also appropriate when the MEP must face into the provider service/bridge domain. An upward-facing MEP sends and receives CFM frames toward the provider network, allowing the maintenance association to span the VPLS service between the provider switches. Together, upward MEP placement and CFM CCM monitoring provide service-level, end-to-end failure detection rather than only physical-link or directly adjacent-neighbor monitoring. Cisco documents CFM as the Ethernet OAM mechanism that uses MEPs, maintenance domains, and CCMs for continuity checking; see [Cisco Ethernet CFM Configuration Guide](#) and [IEEE 802.1ag Connectivity Fault Management](#).

QUESTION NO: 31

Which two factors must be considered for high availability in campus LAN designs to mitigate concerns about unavailability of network resources? (Choose two.)

- A. device resiliency
- B. device type
- C. network type
- D. network resiliency
- E. network size

ANSWER: A D

Explanation:

Device resiliency and network resiliency are the two essential high-availability considerations in a campus LAN. Device resiliency addresses failure tolerance within an individual network device. A campus design should use resilient infrastructure features appropriate to the role and availability target, such as redundant power supplies, redundant supervisors or route processors on supported modular platforms, hot-swappable components, and software capabilities that minimize service disruption during a fault or maintenance event. This reduces the chance that a single hardware or software component failure removes a critical network function.

Network resiliency addresses availability across the topology rather than within one chassis. It is achieved by eliminating single points of failure with diverse physical links and redundant devices, defining appropriate Layer 2 and Layer 3 boundaries, and using routing, gateway, and switching protocols that converge predictably after failures. Together, these dimensions allow the campus to continue forwarding traffic or restore it quickly when a device, link, power source, or forwarding path fails. Cisco's campus design guidance treats resilient hardware and a resilient network architecture as complementary elements of a highly available enterprise campus. See the [Cisco Campus Design Zone](#) and the [Cisco Enterprise Campus Network overview](#).

QUESTION NO: 32

A healthcare customer requested that SNMP traps must be sent over the MPLS Layer 3 VPN service. Which protocol must be enabled?

- A. SNMPv3
- B. Syslog
- C. Syslog TLS
- D. SNMPv2
- E. SSH

ANSWER: A

Explanation:

SNMPv3 must be enabled because it provides the security capabilities needed to send management notifications, including SNMP traps, across an MPLS Layer 3 VPN service. An MPLS L3VPN uses VRF-based traffic separation, which provides logical isolation between VPN customers, but it is not an application-layer security mechanism for the contents of management messages. For a healthcare environment, trap data can expose operationally sensitive information such as device identities, interface states, alarms, and topology details. SNMPv3 protects this traffic through the User-based Security Model (USM). USM supports message authentication to verify the sender and detect message modification, as well as privacy through encryption of the SNMP payload. A trap receiver can therefore validate that notifications originate from an authorized managed device and, when privacy is configured, prevent parties with access to the transport from reading notification contents. Cisco documents SNMPv3 as the secure version of SNMP and supports authentication and privacy configuration for SNMP notifications. The protocol security model is standardized in RFC 3414. See [Cisco SNMP configuration and security guidance](#) and [RFC 3414: User-based Security Model for SNMPv3](#).

QUESTION NO: 33

You are tasked with the design of a high available network. Which two features provide fail closed environments? (Choose two.)

- A. EIGRP
- B. RPVST+
- C. MST
- D. L2MP

ANSWER: B C

Explanation:

RPVST+ and MST provide a fail-closed posture for Layer 2 high-availability designs because they use spanning-tree calculations to prevent redundant Ethernet links from creating forwarding loops. In a stable topology, each protocol deliberately places selected redundant ports into a non-forwarding state. A port is permitted to forward only after the protocol has established that forwarding is consistent with a loop-free tree. This conservative default limits the risk of Layer 2 loops, flooding, and broadcast storms when topology information is incomplete or changes are being processed.

RPVST+ applies rapid spanning-tree behavior separately to VLANs, enabling load distribution across links while retaining loop prevention for each VLAN topology. MST similarly provides loop-free Layer 2 forwarding, but maps VLANs to a smaller set of spanning-tree instances, improving scalability while maintaining controlled blocked and forwarding port roles. During failures and topology reconvergence, both mechanisms recalculate the safe active topology before allowing affected redundant paths to forward. This is the essential fail-closed property: forwarding is withheld where safety cannot yet be established. Cisco documents the STP port-role and state model, including the discarding state used before safe forwarding, in its [Spanning Tree Protocol overview](#) and describes MST operation in the [Cisco MST configuration guide](#).

QUESTION NO: 34

Which CIA triad principle is used by social media platforms to constitute a standard procedure of user IDs and passwords requirements?

- A. integrity
- B. confidentiality
- C. availability
- D. compliance

ANSWER: B

Explanation:

confidentiality is correct because user ID and password requirements establish an authentication process that helps ensure a social-media account, its private messages, profile data, and account-management functions are accessible only by the authorized user. In the CIA triad, confidentiality is the objective of preventing unauthorized disclosure of information. Authentication mechanisms, including passwords and account credentials, are common administrative and technical controls used to support that objective: the platform first verifies an asserted identity, then can limit access to the information and actions associated with that identity. For example, requiring a unique user ID and a sufficiently strong password reduces the likelihood that an unauthorized person can sign in and view direct messages, private posts, account details, or other nonpublic content. Password policies may also require minimum length, complexity, reuse restrictions, or multifactor authentication to make this confidentiality protection more resilient. NIST defines confidentiality as preserving authorized restrictions on information access and disclosure, including protecting personal privacy and proprietary information. Cisco likewise describes confidentiality as ensuring that sensitive data is accessed only by authorized users. See the [NIST definition of confidentiality](#) and Cisco's overview of the [CIA triad in cybersecurity](#).

QUESTION NO: 35

Which two conditions must be met for EIGRP to maintain an alternate loop-free path to a remote network? (Choose two.)

- A. The Reported Distance from a feasible successor is lower than the local Feasible Distance.
- B. The Reported Distance from a successor is higher than the local Feasible Distance.
- C. The feasibility condition does not need to be met.
- D. The Feasible Distance from a successor is lower than the local Reported Distance.
- E. A feasible successor must be present.

ANSWER: A E

Explanation:

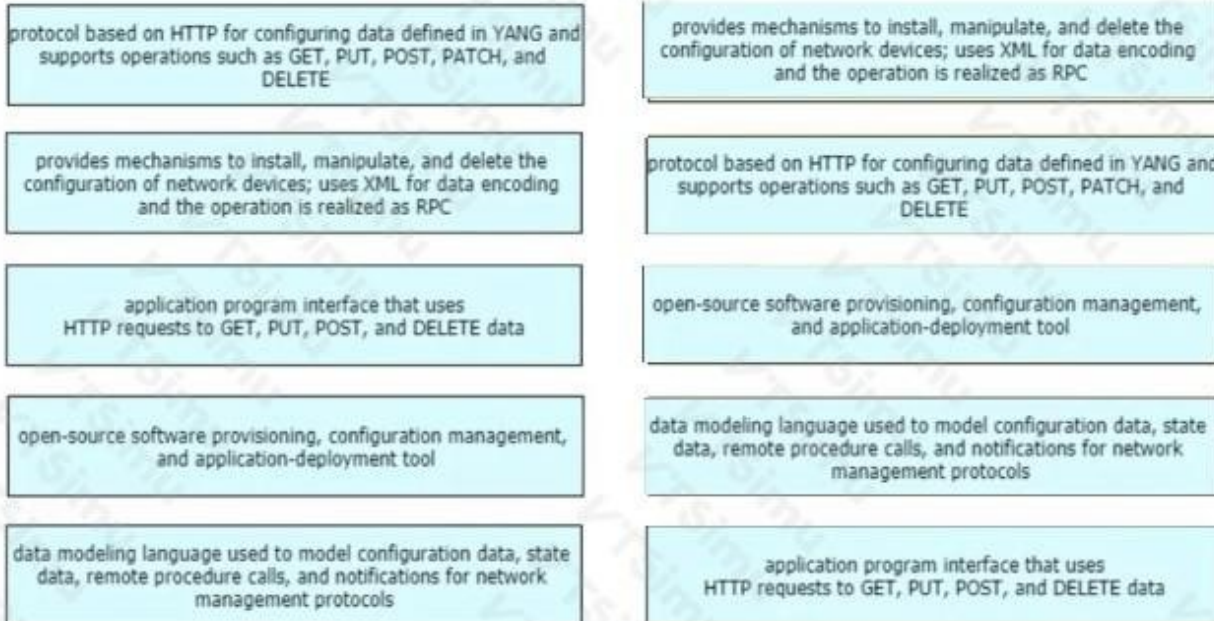
EIGRP's Diffusing Update Algorithm maintains a precomputed alternate path in the topology table as a feasible successor. To qualify, the candidate neighbor's Reported Distance—its advertised metric from that neighbor to the destination—must be strictly lower than the local router's current Feasible Distance for that destination. This is the feasibility condition. It establishes that the neighbor is already closer to the destination than the local router's best known distance, which provides the loop-free guarantee required to install the backup immediately if the successor path fails. The candidate's total metric through that neighbor may be higher than the successor's metric, so it can remain a backup route rather than the selected successor; the essential loop-prevention test is the relationship between Reported Distance and Feasible Distance. A feasible successor must also actually be present in the topology table. When an existing feasible successor is available, EIGRP can promote it without sending queries or waiting for a network-wide recomputation, which is the basis for EIGRP's rapid convergence behavior. Cisco documents the feasibility condition as the requirement that a neighbor's advertised distance be less than the local feasible distance, and describes feasible successors as loop-free backup routes. See [Cisco: EIGRP Feasible Successor and Feasibility Condition](#) and [Cisco: EIGRP Technical Documentation](#).

QUESTION NO: 36 - (DRAG DROP)

Drag and drop the characteristics from the left onto the corresponding network management options on the right.

protocol based on HTTP for configuring data defined in YANG and supports operations such as GET, PUT, POST, PATCH, and DELETE	NETCONF
provides mechanisms to install, manipulate, and delete the configuration of network devices; uses XML for data encoding and the operation is realized as RPC	RESTCONF
application program interface that uses HTTP requests to GET, PUT, POST, and DELETE data	ANSIBLE
open-source software provisioning, configuration management, and application-deployment tool	YANG
data modeling language used to model configuration data, state data, remote procedure calls, and notifications for network management protocols	REST API

ANSWER:



Explanation:

The mapping is correct because these technologies have clearly defined, complementary roles in model-driven network management. NETCONF provides a structured protocol for changing and retrieving device configuration. Its protocol messages are encoded in XML, and its management actions are invoked through remote procedure call operations. This directly identifies the XML and RPC characteristic with NETCONF, as defined by [RFC 6241](#).

RESTCONF provides a RESTful interface to data modeled with YANG. It runs over HTTP and uses standard methods such as GET, PUT, POST, PATCH, and DELETE to retrieve, create, replace, modify, or remove configuration and operational resources. This is the intended HTTP-and-YANG pairing described in [RFC 8040](#).

YANG supplies the data-modeling language that defines configuration data, operational state, remote procedure calls, and notifications in a consistent schema. The scope of that model is established in [RFC 7950](#), making the data-modeling characteristic the proper match for YANG.

Ansible is an open-source automation platform used for provisioning, configuration management, and application deployment. Its automation role and agentless workflow are documented in the [Ansible documentation](#). Finally, a REST API is the general HTTP-based application interface pattern through which resources are accessed and manipulated with request methods such as GET, PUT, POST, and DELETE. The displayed placements therefore correctly distinguish the generic API style from RESTCONF, the specific standardized YANG data-management protocol.

QUESTION NO: 37

If the desire is to connect virtual network functions together to accommodate different types of network service connectivity, what must be deployed?

- A. Bridging
- B. Service Chaining
- C. Linking
- D. Daisy Chaining
- E. Switching

ANSWER: B

Explanation:

Service Chaining must be deployed when traffic needs to traverse a defined sequence of virtual network functions (VNFs) to deliver a network service. In an NFV architecture, individual functions such as virtual firewalls, load balancers, NAT gateways, intrusion-prevention systems, and WAN optimization functions are instantiated separately. Service chaining provides the traffic-steering and policy framework that connects those functions in the required logical order for a given application, tenant, or traffic class. The sequence can be adjusted as service requirements change, without requiring the physical topology to be redesigned. This separation of service logic from the underlying forwarding infrastructure is a core benefit of virtualized network services: operators can compose, automate, scale, and modify end-to-end services through orchestration and policy. The ETSI NFV architectural framework identifies service chaining as a relevant capability for composing NFV network services from virtualized functions and their connectivity requirements. Cisco's NFV materials likewise describe service chaining as steering traffic through an ordered set of network functions. See the [ETSI NFV overview](#) and Cisco's [Network Functions Virtualization overview](#).

QUESTION NO: 38

To provide network resilience organizations need to adopt a holistic approach that includes several key practices and technologies What are two effective ways to enhance network resilience by providing backup and alternative options to maintain network functionality?

- A. scalability
- B. flexibility
- C. recovery
- D. diversity
- E. redundancy

ANSWER: D E

Explanation:

diversity and **redundancy** are the two resilience practices that directly provide backup and alternate means of maintaining connectivity when failures occur. Redundancy supplies additional network components or paths that can assume service if a primary device, link, or service fails. Typical implementations include dual routers or switches, redundant power supplies, paired firewalls, multiple WAN circuits, and dynamic routing or first-hop redundancy protocols that support rapid failover.

Diversity complements redundancy by reducing the chance that the primary and backup resources fail from the same event. A resilient design uses independent failure domains, such as physically separate cable paths, different building entrances, geographically distinct sites, independent carriers, or alternate transport technologies. For example, two circuits are more resilient when they use separate providers and diverse last-mile routes rather than sharing the same conduit or provider infrastructure.

Together, redundancy provides an available standby capability, while diversity protects that capability from common-mode failures. This combination is a core design principle for delivering high availability and preserving network functionality through infrastructure faults or provider outages. Cisco discusses redundant topology and path design in its [Campus High Availability Design Guide](#) and resilient enterprise network design concepts in the [Enterprise Campus Network Design Guide](#).

QUESTION NO: 39

Which two characteristics are associated with 802.1s? (Choose two)

- A. 802.1s supports up to 1024 instances of 802.1
- B. 802.1 s is a Cisco enhancement to 802.1w.
- C. 802.1s provides for faster convergence over 802.1D and PVST+.
- D. CPU and memory requirements are the highest of all spanning-tree STP implementations.
- E. 802.1s maps multiple VLANs to the same spanning-tree instance

ANSWER: C E

Explanation:

IEEE 802.1s, commonly called Multiple Spanning Tree Protocol (MSTP), combines the rapid reconvergence behavior of RSTP with the ability to use a reduced number of spanning-tree instances. Therefore, “802.1s provides for faster convergence over 802.1D and PVST+” is a characteristic of MSTP. Within an MST region, the protocol uses the rapid port-state transition and topology-change mechanisms derived from IEEE 802.1w/RSTP. This enables significantly faster recovery from many Layer 2 topology changes than traditional IEEE 802.1D operation and classic PVST+ behavior.

“802.1s maps multiple VLANs to the same spanning-tree instance” is also fundamental to MSTP. Administrators associate VLANs with MST instances, allowing many VLANs to share one tree rather than requiring a separate tree calculation for every VLAN. This reduces control-plane processing and memory consumption while retaining the ability to distribute VLAN forwarding paths across multiple instances. MST instance zero, the Internal Spanning Tree, provides connectivity within the region, while additional MST instances can be configured for mapped VLAN groups. Cisco documents MST as an IEEE standard protocol that maps VLANs to spanning-tree instances and incorporates rapid convergence behavior. See Cisco’s [Spanning Tree Protocol overview](#) and [MSTP configuration guide](#).

QUESTION NO: 40 - (SIMULATION)

Indicate the nature of automation and orchestration tasks by dragging the tasks on the left to the corresponding category on the right in no particular order.

Reset the OSPF process on three routers.	Automation	task
Send the configuration tasks to the hardware device.		task
Create a new VLAN on all your switches.		task
Translate request into configuration tasks for network devices to execute.	Orchestration	task
Monitor progress of the tasks sent to the hardware device.		task
Create a new loopback interface on five routers.		task

ANSWER: See the explanation for the answer

Explanation:

Automation consists of the repeatable device/configuration actions:

Orchestration coordinates the workflow that turns a request into actions, deploys them, and tracks their completion:

Automation performs defined, repeatable operations; orchestration manages the end-to-end workflow and coordination of those operations.

QUESTION NO: 41

You are designing a network for a branch office. In order to improve convergence time, you are required to use the BFD feature Which four routing protocols can you use to facilitate this? (Choose four.)

A. IS-IS

- B. static
- C. RIP
- D. EIGRP
- E. BGP

ANSWER: A B D E

Explanation:

Bidirectional Forwarding Detection (BFD) is designed to provide rapid, protocol-independent detection of a forwarding-path failure between adjacent devices. Rather than relying solely on comparatively long hello, hold, or keepalive timers, BFD communicates session failure to the route-processing client, allowing the control plane to withdraw affected routes and converge quickly. This is particularly valuable for a branch design, where loss of a WAN circuit, provider path, or next-hop device must trigger fast failover to an alternate connection.

IS-IS, EIGRP, and BGP can each register with BFD and use the BFD session state to promptly tear down an adjacency or peering when bidirectional forwarding is lost. This permits fast convergence without requiring excessively aggressive native protocol timers. Static routes can also use BFD tracking. When the associated BFD session fails, the static route is removed from the routing table, allowing a backup static route or a dynamically learned alternative path to become active. This makes static-route BFD especially useful for branch default-route and dual-uplink designs. Cisco documents BFD as a common failure-detection mechanism for supported routing-protocol clients and static routing use cases. See the [Cisco BFD overview](#) and the [Cisco IOS XE BFD Configuration Guide](#).

QUESTION NO: 42

IPFIX data collection via standalone IPFIX probes is an alternative to flow collection from routers and switches. Which use case is suitable for using IPFIX probes?

- A. performance monitoring
- B. security
- C. observation of critical links
- D. capacity planning

ANSWER: C

Explanation:

Standalone IPFIX probes are especially suitable for **observation of critical links**. A probe receives a copy of packets from a network TAP or switch SPAN/RSPAN/ERSPAN session and acts as the IPFIX metering process at that observation point. This model is valuable when detailed, independent visibility is required on a limited number of high-value paths, such as an Internet edge, a data-center interconnect, a WAN handoff, or a regulated application segment.

Because the probe is dedicated to traffic observation rather than packet forwarding, it can collect and export flow telemetry without consuming forwarding-device resources or relying on the router or switch to support the necessary flow-export features and scale. It also allows the designer to place the observation point precisely where traffic must be inspected, including links that involve devices unable to export flow records. The IPFIX architecture explicitly defines observation points and metering processes, which can be implemented by a dedicated probe observing packets on a selected link. This makes focused monitoring of strategically important links the appropriate deployment use case. See [RFC 7011: IPFIX Protocol Specification](#) and [Cisco NetFlow and flow-monitoring guidance](#).

QUESTION NO: 43

An MPLS service provider is offering a standard EoMPLS-based VPLS service to Customer

A. providing Layer 2 connectivity between a central site and approximately 100 remote sites. Customer A wants to use the VPLS network to carry its internal multicast video feeds which are sourced at the central site and consist of 20 groups at Mbps each. Which service provider recommendation offers the most scalability?

EoMPLS-based VPLS can carry multicast traffic in a scalable manner

B. Use a mesh of GRE tunnels to carry the streams between sites

C. Enable snooping mechanisms on the provider PE routers.

D. Replace VPLS with a Layer 3 MVPN solution to carry the streams between sites

ANSWER: D

Explanation:

Replace VPLS with a Layer 3 MVPN solution to carry the streams between sites is the most scalable recommendation for this hub-and-spoke multicast video requirement. Multicast VPN extends Layer 3 VPNs with multicast-aware customer and provider control-plane mechanisms, allowing the provider to build multicast distribution trees for only the VPN sites that have active receivers. Customer multicast routing information is maintained in the relevant VRF, while the MPLS provider backbone can transport selected multicast flows using technologies such as multipoint LDP, RSVP-TE point-to-multipoint trees, or BGP-based MVPN signaling, depending on the design and platform capabilities.

This architecture is well suited to a central source distributing many high-bandwidth groups to a large number of remote locations. It separates customer multicast routing from provider transport, supports receiver-driven tree construction, and avoids treating all video streams as generic Layer 2 flooded traffic across a multipoint Ethernet service. The result is more efficient use of backbone capacity and a control plane designed to manage multicast membership and replication at VPN scale. Cisco documents MVPN operation and configuration in the [Cisco IOS XE Multicast VPN Configuration Guide](#) and describes multicast VPN design concepts in its [MVPN Technical Overview](#).

QUESTION NO: 44

When designing a WAN that will be carrying real-time traffic, what are two important reasons to consider serialization delay? (Choose two)

A. Serialization delays are invariable because they depend only on the line rate of the interface

B. Serialization delays are variable because they depend on the line rate of the interface and on the type of the packet being serialized.

C. Serialization delay is the time required to transmit the packet on the physical media.

D. Serialization delays are variable because they depend only on the size of the packet being serialized

E. Serialization delay depends not only on the line rate of the interface but also on the size of the packet

ANSWER: C E

Explanation:

Serialization delay is the time needed to place all bits of a packet onto a physical interface. In a WAN design, this is an important component of per-hop delay because an interface cannot transmit the next packet until the current packet has been serialized. For real-time voice and video, calculating this delay helps determine whether the network can meet the end-to-end latency and jitter budget.

The serialization-delay calculation depends on both packet size and the interface transmission rate: $\text{serialization delay} = \text{packet size in bits} / \text{link speed in bits per second}$. Consequently, a larger frame takes longer to transmit, while a lower-speed WAN circuit increases the transmission time for every frame. This relationship is particularly significant on slow-speed access links, where a delay-sensitive packet can be held behind a larger packet already being sent. Designers use this knowledge when evaluating QoS requirements and whether mechanisms such as link fragmentation and interleaving, traffic shaping, or a higher-speed circuit are needed to protect real-time traffic.

Cisco describes serialization as a delay component associated with transmitting a packet onto an interface, and its QoS guidance addresses the effects of packet size and bandwidth on delay-sensitive traffic. See [Cisco QoS Overview and Troubleshooting](#) and [Cisco QoS Concepts](#).

QUESTION NO: 45

When planning their cloud migration journey, what is crucial for virtually all organizations to perform?

- A. SASE framework deployment
- B. Optimizing the WAN environment
- C. Assessment of current infrastructure
- D. RPO and RTO calculations duration planning

ANSWER: C

Explanation:

Assessment of current infrastructure is crucial because a cloud migration must begin with a reliable understanding of the organization's existing environment and business requirements. The assessment inventories applications, servers, data stores, network connectivity, identities, security controls, operational dependencies, and regulatory obligations. It also establishes each workload's resource consumption, criticality, performance requirements, availability needs, and suitability for a particular migration approach, such as rehosting, replatforming, refactoring, retaining, or retiring.

This discovery work identifies hidden application dependencies and data flows that could otherwise cause outages, security gaps, unexpected latency, or migration sequencing problems. It gives architects the evidence needed to build a phased migration roadmap, size cloud connectivity and landing-zone services, estimate costs, define governance, and prioritize workloads according to business value and risk. The assessment is not merely a technical inventory: stakeholder requirements, operating-model readiness, and compliance constraints must also be captured so that the target design supports the organization after migration. Cisco describes cloud migration as a journey requiring an integrated strategy across applications, infrastructure, security, and operations; its cloud resources are available at [Cisco Cloud Solutions](#). Cisco's guidance on workload and infrastructure visibility is also reflected in its [Cisco Intersight Workload Optimizer](#) resources.

QUESTION NO: 46

: 488

The modularity built into the architecture allows flexibility in network design and facilitates implementation and troubleshooting Which solution is difficult to implement manage and troubleshoot especially for large networks?

- A. functional boundaries
- B. logical core layers
- C. distribution network
- D. hierarchical model
- E. flat network design

ANSWER: E

Explanation:

flat network design is correct because it lacks the modular functional separation used in a hierarchical campus architecture. In a flat design, network devices and services are generally placed in one broadly interconnected layer rather than being organized into distinct access, distribution, and core functions. As the network grows, this makes operational changes harder to contain and complicates fault isolation: a problem can affect a larger portion of the network, and engineers have fewer

clearly defined boundaries at which to apply policy, summarize routes, enforce security, or narrow a troubleshooting investigation.

Cisco's enterprise campus approach uses hierarchical modularity to make networks more scalable, resilient, and manageable. The access, distribution, and core roles create structured control points and reduce the scope of failures and changes. A flat network does not provide those architectural benefits, so it becomes increasingly difficult to deploy consistently, operate efficiently, and troubleshoot as the number of endpoints, devices, and services increases. See Cisco's [Campus Networking overview](#) and the [Enterprise Networking solutions](#) documentation.

QUESTION NO: 47

You want to split an Ethernet domain in two.

Which parameter must be unique in this design to keep the two domains separated?

- A. VTP domain
- B. VTP password
- C. STP type
- D. VLAN ID

ANSWER: D

Explanation:

VLAN ID is the required unique parameter because a VLAN defines a separate Layer 2 broadcast domain on Ethernet switches. Ports assigned to different VLAN IDs are logically separated even when they use the same physical switching infrastructure. Switches use the VLAN identifier to determine the forwarding and flooding scope for Ethernet frames; broadcasts, unknown unicasts, and multicast traffic are contained within the applicable VLAN. On an IEEE 802.1Q trunk, the VLAN ID is carried in the frame tag, allowing multiple independent Layer 2 domains to traverse one physical link while remaining isolated from one another. Therefore, assigning distinct VLAN IDs is the fundamental design mechanism for splitting one Ethernet switching environment into two separate domains. Inter-VLAN communication, if needed, then requires a Layer 3 routing function and an explicit routing policy. Cisco describes VLANs as logical broadcast domains and explains that VLAN tagging identifies the VLAN to which traffic belongs. See [Cisco's VLAN design overview](#) and the [Cisco VLAN trunk configuration guide](#).

QUESTION NO: 48

A large enterprise is planning a new WAN connection to headquarters. The current dual-homed setup with static routing is not providing consistent resiliency. Users complain when one specific link fails, while failure of the other causes no issues. The organization wants to improve resiliency and ROI.

Which solution should be recommended?

- A. Implement granular quality of service on the links
- B. Procure additional bandwidth
- C. Use dynamic routing toward the WAN
- D. Add an additional link to the WAN

ANSWER: C

Explanation:

Use dynamic routing toward the WAN is the appropriate recommendation because it addresses the underlying cause: static routes do not inherently provide adaptive, end-to-end path selection when a WAN path becomes unavailable or is only partially reachable. A dynamic routing protocol can exchange reachability information between the enterprise edge and

headquarters, select paths based on defined routing policy and metrics, and withdraw or reconverge routes when a failure is detected. This enables traffic to move to the remaining viable connection without the manual route changes or brittle tracking logic that commonly accompany static-routing designs.

The design should also use appropriate failure detection and convergence tuning so that loss of a link, next hop, or end-to-end forwarding capability is recognized promptly. For example, Bidirectional Forwarding Detection can be associated with supported routing protocols to detect forwarding failures independently of slower protocol timers. Dynamic routing therefore improves availability while making better use of the existing dual-homed WAN investment, supporting the resiliency and ROI objectives without requiring an additional circuit as the primary remedy. Cisco documents dynamic-routing operation and route exchange in its [OSPF overview](#) and explains rapid failure detection in its [BFD configuration guide](#).

QUESTION NO: 49

What are two advantages of controller-based networks versus traditional networks? (Choose two.)

- A. the ability to have forwarding tables at each device
- B. more flexible configuration per device
- C. more consistent device configuration
- D. programmatic APIs that are available per device
- E. the ability to configure the features for the network rather than per device

ANSWER: C E

Explanation:

Controller-based networking provides centralized management, automation, and policy definition, allowing network intent to be applied uniformly across the infrastructure. This produces **more consistent device configuration** because the controller can use common templates, standardized policy, and automated provisioning to reduce manual variation and configuration drift. Operators define the intended network behavior once, and the controller translates and deploys the required device-level configuration in a consistent manner.

It also provides **the ability to configure the features for the network rather than per device**. Network-wide capabilities such as segmentation, quality of service, security policy, and access controls can be defined centrally according to business or operational intent. The controller then orchestrates implementation across the relevant devices and continuously helps maintain the desired state. This shifts operations away from individually configuring each router and switch, improving scale, repeatability, and the speed of controlled changes.

Cisco describes software-defined networking as separating centralized control and policy from the underlying infrastructure, enabling automation and consistent policy deployment. Cisco Catalyst Center similarly provides centralized automation and policy-based management for enterprise networks. See [Cisco Software-Defined Networking](#) and [Cisco Catalyst Center](#).

QUESTION NO: 50

A company plans to use BFD between its routers to detect a connectivity problem inside the switched network. An IPS is transparently installed between the switches. Which packets should the IPS forward for BFD to work under all circumstances?

- A. Fragmented packet with the do-not-fragment bit set
- B. IP packets with broadcast IP source addresses
- C. IP packets with the multicast IP source address
- D. IP packet with the multicast IP destination address
- E. IP packets with identical source and destination IP addresses
- F. IP packets with the destination IP address 0.0.0.0.

ANSWER: D

Explanation:

IP packet with the multicast IP destination address is correct because single-hop BFD uses the dedicated IPv4 multicast destination address 224.0.0.184 for BFD Control packets, particularly while a session is being established or re-established before the peer discriminator is known. Therefore, an inline transparent IPS must pass this multicast-destination traffic without filtering it, even if its normal security policy treats multicast as exceptional traffic. If these packets do not traverse the switched path, adjacent routers cannot reliably initiate BFD or recover BFD operation after a session reset.

BFD is designed to provide rapid, protocol-independent forwarding-path failure detection. In a single-hop deployment, the packets are also constrained by single-hop validation requirements, including a TTL/Hop Limit of 255, so they remain local to the directly connected forwarding domain. The IPS is part of that domain from the routers' perspective; consequently, it must transparently forward the BFD multicast packets rather than block, proxy, or alter them. Cisco's BFD documentation identifies 224.0.0.184 as the IPv4 multicast address used by BFD, and RFC 5881 specifies the single-hop BFD packet format and use of this destination address. See [RFC 5881](#) and [Cisco IOS XE BFD Configuration Guide](#).

QUESTION NO: 51

Which two protocols are used by SDN controllers to communicate with switches and routers? (Choose two)

- A. OpenFlash
- B. OpenFlow
- C. NetFlash
- D. Open vSwitch Database
- E. NetFlow

ANSWER: B D

Explanation:

OpenFlow and Open vSwitch Database are southbound protocols that an SDN controller can use to communicate with forwarding devices. OpenFlow defines a standardized controller-to-datapath interface through which the controller can program forwarding behavior, including adding, modifying, and removing flow entries. This supports the fundamental SDN model in which a logically centralized control plane directs how switches forward traffic. The Open Networking Foundation describes OpenFlow as a protocol enabling an external controller to access and control a switch's forwarding plane.

Open vSwitch Database (OVSDB) complements flow programming by providing a management protocol for configuring and monitoring Open vSwitch instances. A controller can use OVSDB to create and manage bridges and ports, configure tunnels, apply VLAN-related settings, and obtain configuration or operational state. In deployments using Open vSwitch, it is common for OVSDB to handle device configuration while OpenFlow installs forwarding policies. The OVSDB protocol is specified in RFC 7047, which defines the database-management interface used to communicate with Open vSwitch. See the [Open Networking Foundation's OpenFlow overview](#) and [RFC 7047: Open vSwitch Database Management Protocol](#).

QUESTION NO: 52

Which two design options are available to dynamically discover the RP in an IPv6 multicast network? (Choose two)

- A. Embedded RP
- B. MSDP
- C. BSR
- D. Auto-RP
- E. MLD

ANSWER: A C**Explanation:**

Embedded RP and BSR are the two available mechanisms for dynamically identifying a PIM rendezvous point in an IPv6 multicast design. Embedded RP uses an IPv6 multicast group address format that carries the rendezvous-point address within the group address itself. Routers can derive the rendezvous point directly from the group address, avoiding a separate rendezvous-point mapping distribution process for multicast groups designed with that addressing model. This approach is defined for IPv6 multicast addressing and is particularly useful when multicast addressing can be planned to encode rendezvous-point information.

BSR, or Bootstrap Router, provides dynamic rendezvous-point discovery by distributing rendezvous-point and group-range information throughout the PIM domain. Candidate rendezvous points advertise their availability to the bootstrap router, and bootstrap messages disseminate the resulting mapping information to PIM routers. This allows rendezvous-point assignments to be learned dynamically and supports redundant candidate rendezvous points as part of a scalable domain design. Cisco documents IPv6 PIM support for both bootstrap rendezvous-point discovery and embedded rendezvous points. See [Cisco IPv6 PIM Configuration Guide](#) and [RFC 3956: Embedded-RP Addressing](#).

QUESTION NO: 53

Company XYZ, a global content provider, owns data centers on different continents. Their data center design involves a standard three-layer design with a Layer 3-only core. VRRP is used as the FHRP. They require VLAN extension across access switches in all data centers and plan to purchase a Layer 2 interconnection between two of their data centers in Europe. In the absence of other business or technical constraints, which termination point is optimal for the Layer 2 interconnection?

- A. At the core layer, to offer the possibility to isolate STP domains
- B. At the access layer because the STP root bridge does not need to align with the VRRP active node
- C. At the core layer because all external connections must terminate there for security reasons
- D. At the aggregation layer because it is the Layer 2 to Layer 3 demarcation point

ANSWER: D**Explanation:**

At the aggregation layer because it is the Layer 2 to Layer 3 demarcation point is correct. In a conventional three-tier data-center topology, the aggregation layer is the natural boundary between the Layer 2 access domain and the routed core. Terminating the inter-data-center Layer 2 service there extends the required VLANs while preserving the core as a Layer 3-only transport domain. This maintains the intended hierarchy: access switches connect to an aggregation pair that provides VLAN forwarding, first-hop gateway services such as VRRP, and policy enforcement, while the core provides scalable routed connectivity between aggregation blocks.

This placement also gives the design a clear, controllable point for the Layer 2 domain that traverses the interconnection. The aggregation devices can consistently carry the required VLANs, participate in the relevant spanning-tree topology, and provide the associated gateway functions without introducing bridged forwarding into the routed core. Keeping Layer 2 extension at this demarcation supports modularity and limits the scope of the extended broadcast and control-plane domain to the part of the design that is intended to provide Layer 2 services. Cisco's hierarchical data-center design guidance describes the aggregation layer as the point that aggregates access-layer connectivity and provides the Layer 2/Layer 3 boundary; see [Cisco Data Center Infrastructure Design Guide](#) and [Cisco Data Center Networking](#).

QUESTION NO: 54

Scalability is a desirable attribute of a network, system, or process. Poor scalability can result in poor system performance, necessitating the reengineering or duplication of systems. Load scalability is the ability of a system to perform gracefully as traffic increases. Which two problems can occur due to poor load scalability design? (Choose two.)

- A. cannot fully take advantage of parallelism

- B. algorithmically intolerable
- C. limited size of a data structure
- D. repeatedly engaging in wasteful activity
- E. redundant message logging

ANSWER: A D

Explanation:

cannot fully take advantage of parallelism and **repeatedly engaging in wasteful activity** are load-scalability problems because both cause the work required per unit of traffic to grow inefficiently as demand rises. A design fails to fully exploit parallelism when processing remains serialized around a shared resource, control point, or synchronization boundary. Even if additional compute, links, or devices are introduced, throughput does not increase proportionally because the system cannot distribute independent work effectively. In a networked service, this can appear as a single control-plane process, database operation, or packet-processing stage becoming the limiting factor.

Repeatedly engaging in wasteful activity also degrades graceful behavior under load. Examples include recalculating information that could be cached, repeatedly polling or flooding for state that is already known, excessive retries, or processing duplicate work. As traffic grows, this unnecessary work consumes CPU, memory, bandwidth, and queue capacity, leaving fewer resources for useful forwarding or service operations. Sound scalable design therefore seeks to partition work safely, minimize shared bottlenecks, cache reusable results where appropriate, and avoid duplicate processing. Cisco design guidance emphasizes identifying bottlenecks and designing systems that can expand capacity while maintaining predictable operation: [Cisco Enterprise Network Design](#). General distributed-systems scalability principles, including parallel execution and avoiding unnecessary coordination, are also summarized in [Google SRE: Handling Overload](#).

QUESTION NO: 55

A business requirement stating that failure of WAN access for dual circuits into an MPLS provider for a Data Centre cannot happen due to related service credits that would need to be paid has led to diversely routed circuits to different points of presence on the providers network? What should a network designer also consider as part of the requirement?

- A. Provision of an additional MPLS provider
- B. Out of band access to the MPLS routers
- C. Ensuring all related remote branches are dual homed to the MPLS network
- D. Dual PSUs & Supervisors on each MPLS router

ANSWER: A

Explanation:

Provision of an additional MPLS provider is the appropriate additional consideration because path diversity to separate points of presence within one service-provider network does not eliminate the service-provider failure domain. Although the access circuits may take physically diverse routes and terminate on different provider edge locations, both still depend on the same carrier's backbone, operational processes, control-plane design, maintenance practices, and potentially shared transport or platform infrastructure. A large-scale provider incident can therefore affect both circuits at once.

For a requirement in which loss of data-center WAN access is unacceptable, the design must extend redundancy across independent administrative and operational domains. Procuring connectivity from a second MPLS carrier provides carrier diversity and substantially reduces common-mode risk. The data-center edge should then be designed for resilient dual-provider routing, with deliberate path-selection and failover policies, capacity sufficient to carry traffic after a provider failure, and regular testing of the failover procedure. The business requirement should also define availability targets, restoration expectations, and the precise demarcation of provider responsibilities. Cisco WAN design guidance emphasizes designing resiliency around independent failure domains rather than relying only on duplicate links within a shared domain. See Cisco's [WAN Design Zone](#) and [Enterprise Networking solutions](#).

QUESTION NO: 56

Which relationship between IBGP and the underlying physical topology is true?

- A. iBGP full mesh requirement does not dictate any specific network topology.
- B. iBGP can work only on a ring network topology with a link-state protocol like OSPF or IS-IS
- C. iBGP full mesh requires an underlying fully meshed network topology.
- D. iBGP does not work on a ring network topology even with an underlying IGP.

ANSWER: A

Explanation:

iBGP full mesh requirement does not dictate any specific network topology. The full-mesh rule is a logical BGP-control-plane requirement: absent route reflectors or confederations, each iBGP speaker in the autonomous system must maintain an iBGP session with every other iBGP speaker so that BGP-learned routes can be distributed correctly. It does not mean that every router must have a physical link to every other router. An iBGP session uses TCP port 179 and can be established over multiple routed hops, provided that each peer address has stable bidirectional IP reachability. That reachability is commonly supplied by an underlay IGP such as OSPF or IS-IS, but static routes or another suitable routing design can also provide it. Therefore, a ring, hub-and-spoke, partial mesh, or other physical design can support a logical iBGP full mesh. At larger scale, route reflectors or confederations can reduce the number of required iBGP sessions; these mechanisms alter the logical peering design without imposing a particular physical topology. Cisco describes iBGP as requiring a full mesh between iBGP peers unless route reflection or confederations are used, while the transport network provides the IP reachability on which those peerings depend. See [Cisco BGP Case Studies](#) and [Cisco BGP Best Practices](#).

QUESTION NO: 57

Enterprise XYZ wants to implement fast convergence on their network and optimize timers for OSPF. However, they also want to prevent excess flooding of LSAs if there is a constantly flapping link on the network. Which timers can help prevent excess flooding of LSAs for OSPF?

- A. OSPF propagation timers
- B. OSPF throttling timers
- C. OSPF delay timers
- D. OSPF flooding timers

ANSWER: B

Explanation:

OSPF throttling timers are correct because they provide a controlled backoff mechanism for control-plane activity during rapid topology changes, such as an unstable or repeatedly flapping link. In particular, LSA-generation throttling controls how quickly a router can originate successive updated LSAs. Rather than immediately generating a new LSA for every closely spaced event, the router applies configured initial, hold, and maximum delays. This reduces repeated LSA origination and the resulting flooding while retaining a short initial response time for an isolated failure. OSPF SPF throttling is typically configured alongside LSA throttling to pace the subsequent shortest-path recalculations, helping preserve CPU capacity and maintain stable convergence behavior at scale. The design objective is not to suppress valid topology information permanently; it is to dampen bursts of control-plane work caused by instability and allow the delay to increase when events continue. Cisco documents these mechanisms through the OSPF `timers lsa-group-pacing` and LSA/SPF timing features, while the OSPF specification defines the flooding and LSA-origination behavior that these operational controls optimize. See Cisco's [OSPF command reference](#) and [RFC 2328, OSPF Version 2](#).

QUESTION NO: 58

You are performing a BGP design review for a service provider that offers MPLS-based services to their end customers. The network is comprised of several PE routers that run iBGP with a pair of route reflectors for all BGP address families. Which two options about the use of Constrained Route Distribution for BGP/MPLS VPNs are true? (Choose two.)

- A. The RRs do not need to advertise any route target filter toward the PE routers.
- B. The RR must advertise the default route target filter toward the PE routers.
- C. Both PE and RR routers must support this feature.
- D. This feature must be enabled on all devices in the network at the same time.

ANSWER: B C

Explanation:

Constrained Route Distribution, commonly called RT-Constrain, improves the scalability of a route-reflector-based BGP/MPLS VPN design. It uses the Route Target membership address family to communicate which Route Targets a PE needs to import. The route reflector can then use that membership information to distribute VPN routes only to PE routers with an interest in the associated Route Targets, rather than reflecting the complete VPN route set to every PE. This reduces unnecessary VPNv4 and VPNv6 control-plane updates, BGP memory consumption, and processing load.

“The RR must advertise the default route target filter toward the PE routers” is correct because the default Route Target filter enables PE routers to advertise their Route Target membership information. This allows the route reflector to construct the membership state required to constrain subsequent VPN-route advertisement. “Both PE and RR routers must support this feature” is also correct: PE routers originate Route Target membership information, while route reflectors receive and apply that information when deciding where to reflect VPN routes. The feature therefore requires compatible Route Target Constraint capability and address-family operation on both ends of the relevant BGP sessions. See [RFC 4684](#) and Cisco’s [RT-Constrain overview](#).

QUESTION NO: 59

A mega store plans to expand its business into the online world, and wants to operate using the highest possible security standards to prove to their customers that they take handling of their payment information seriously. Only TLS v1.3 will be allowed on their websites. Which type of SSL certificate will emphasize their commitment to enforcing high security standards and minimize risk of spoofing?

- A. DV SSL certificate
- B. PV SSL certificate
- C. OV SSL certificate
- D. EV SSL certificate

ANSWER: D

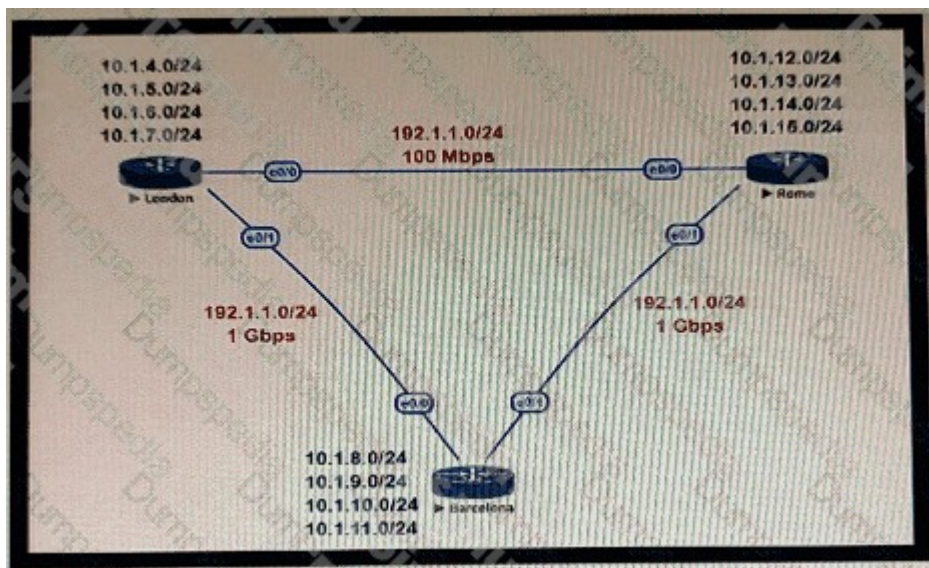
Explanation:

EV SSL certificate is the appropriate choice because Extended Validation certificates require the certificate authority to perform the most rigorous level of organization-identity validation available for publicly trusted TLS certificates. This includes verifying the legal existence and identity of the organization, its physical location, operational status, and the authority of the individual requesting the certificate. That validation creates a substantially stronger binding between the website’s cryptographic identity and the real-world business, helping reduce the risk that an attacker can obtain a lookalike certificate while falsely claiming to represent the retailer.

TLS 1.3 provides modern transport protections, including strong encryption and an improved handshake, but it does not itself establish that the site operator is the legitimate company customers expect. EV validation addresses that identity-assurance requirement, complementing TLS 1.3 for a retailer handling payment information. The validation requirements and issuance controls for EV certificates are defined by the CA/Browser Forum’s [Extended Validation Guidelines](#). TLS 1.3 is standardized in [RFC 8446](#).

QUESTION NO: 60

Refer to the exhibit.



This network is running OSPF as the routing protocol. The internal networks are being advertised in OSPF. London and Rome are using the direct link to reach each other although the transfer rates are better via Barcelona. Which OSPF design change allows OSPF to calculate the proper costs?

- A. Change the OSPF reference bandwidth to accommodate faster links.
- B. Filter the routes on the link between London and Rome.
- C. Change the interface bandwidth on all the links.
- D. Implement OSPF summarization to fix the issue.

ANSWER: A

Explanation:

Change the OSPF reference bandwidth to accommodate faster links. Cisco OSPF derives the default interface cost from the configured reference bandwidth divided by the interface bandwidth. The default reference bandwidth is 100 Mb/s, so interfaces operating at 100 Mb/s or faster can all receive the minimum integer cost of 1. This removes the metric distinction that OSPF needs to prefer a higher-capacity multihop path over a lower-capacity direct path when both paths otherwise have comparable cumulative costs.

Increasing the reference bandwidth to a value that represents the fastest links in the design lets OSPF assign meaningful, differentiated costs to Gigabit and higher-speed interfaces. For example, a reference bandwidth of 100,000 Mb/s supports sensible metric differentiation through 100-Gb/s links. The same `auto-cost reference-bandwidth` value must be configured consistently on every OSPF router in the routing domain, because inconsistent reference values can cause routers to calculate different shortest paths. OSPF then selects paths based on the summed outbound interface costs, allowing the route through Barcelona to be evaluated according to the actual relative link capacities. Cisco documents the reference-bandwidth setting in its [OSPF configuration guide](#); the OSPF metric and shortest-path calculation are defined in [RFC 2328](#).

QUESTION NO: 61

An engineer is designing a DMVPN network where OSPF has been chosen as the routing protocol. A spoke-to-spoke data propagation model must be set up. Which two design considerations must be taken into account? (Choose two)

- A. Configure all the sites as network type broadcast.
- B. The network type on all sites should be point-to-multipoint.

- C. The network type should be point-to-multipoint for the hub and point-to-point for the spokes.
- D. The hub should be set as the DR by specifying the priority to 255.
- E. The hub should be the DR by changing the priority of the spokes to 0.

ANSWER: A E

Explanation:

“Configure all the sites as network type broadcast” is appropriate for an OSPF DMVPN design that supports direct spoke-to-spoke communication. A multipoint GRE tunnel can carry OSPF multicast traffic, and using the broadcast network type places the hub and spokes in one OSPF broadcast segment. This model allows the hub to form the required full adjacencies while spokes maintain the two-way neighbor relationship needed on a broadcast multiaccess network. OSPF route calculation can then select a remote spoke as the appropriate next hop, while DMVPN/NHRP supplies the dynamic tunnel mapping required to build the direct spoke-to-spoke data path.

“The hub should be the DR by changing the priority of the spokes to 0” is also required to keep the OSPF control-plane design stable and scalable. Setting every spoke to OSPF priority zero makes each spoke ineligible for DR and BDR election. The hub is therefore the sole eligible DR, avoiding an unintended spoke becoming DR and ensuring that the hub maintains the full OSPF adjacencies. This is the conventional OSPF broadcast-mode design for a hub-and-spoke DMVPN topology. Cisco describes the spoke-to-spoke DMVPN forwarding behavior at [Cisco DMVPN Spoke-to-Spoke Communication](#) and documents OSPF network-type behavior at [Cisco OSPF Network Types](#).

QUESTION NO: 62

Agile and Waterfall are two popular methods for organizing projects. What describes any Agile network design development process?

- A. working design over comprehensive documentation
- B. contract negotiation over customer collaboration
- C. following a plan over responding to change
- D. processes and tools over individuals and interactions over time

ANSWER: A

Explanation:

“working design over comprehensive documentation” best expresses an Agile approach to network design development. It adapts the Agile Manifesto value of favoring a working outcome over exhaustive documentation to the networking domain. Rather than treating a large, detailed design document as the primary measure of progress, an Agile team delivers small, usable, and verifiable design increments. These can include a lab-validated topology, a tested routing or security-policy pattern, an automation proof of concept, or a pilot deployment that demonstrates stated acceptance criteria.

This emphasis supports short feedback cycles with architecture, operations, security, application, and business stakeholders. Feedback from testing and early implementation can be incorporated into subsequent design iterations, helping the team validate technical assumptions, reveal operational constraints, and reduce delivery risk before broad rollout. Documentation remains important for implementation, operations, governance, and knowledge transfer; Agile simply focuses documentation on what is useful and current while prioritizing evidence that the proposed design works. This principle is consistent with the Agile Manifesto, which states that there is value in documentation but places greater value on a working result, and with iterative delivery practices described by Cisco for network automation and DevNet workflows. See the [Agile Manifesto](#) and [Cisco DevNet learning resources](#).

QUESTION NO: 63

An attacker exploits application flaws to obtain data and credentials. What is the next step after application discovery in Zero Trust networking?

- A. Establish visibility and behavior modeling
- B. Enforce policies and microsegmentation
- C. Assess real-time security health
- D. Ensure trustworthiness of systems

ANSWER: A

Explanation:

Establish visibility and behavior modeling is the appropriate next step because application discovery identifies the applications, services, dependencies, and communications that exist in the environment, but discovery alone does not establish what normal or authorized use looks like. Zero Trust design depends on observing application-to-application, user-to-application, and workload-to-workload flows; identifying identities, devices, and resources involved; and baselining expected behavior. This context enables security teams to recognize anomalous access patterns, understand application dependencies before introducing controls, and create precise least-privilege access policies. Behavior modeling also supplies the evidence needed for continuous evaluation of access requests and risk, which are core Zero Trust principles. Once normal traffic patterns and resource relationships are visible, segmentation and enforcement policies can be designed to protect credentials and sensitive data without unintentionally disrupting legitimate application communications. Cisco describes Zero Trust as an approach that verifies access continuously and applies controls based on context, identity, and policy: [Cisco Zero Trust](#). NIST likewise emphasizes continuous monitoring and evaluation of subjects, devices, and enterprise resources as part of a Zero Trust architecture: [NIST SP 800-207](#).

QUESTION NO: 64

During a pre-sales meeting with a potential customer, the customer CTO asks a question about advantages of controller-based networks versus a traditional network. What are two advantages to mention? (Choose two)

- A. Per device forwarding tables
- B. Programmatic APIs available per device
- C. Abstraction of individual network devices
- D. Distributed control plane
- E. Consistent device configuration

ANSWER: C E

Explanation:

Abstraction of individual network devices is a primary benefit of a controller-based architecture. Rather than requiring operators and applications to account for the specific CLI, capabilities, and configuration state of every switch, router, or access point, the controller can present an intent- or policy-oriented view of the network. This allows the design to express required connectivity, segmentation, security, and service outcomes while the controller translates that intent into device-level implementation. The abstraction makes operations and automation more scalable as the network grows or hardware is refreshed.

Consistent device configuration is also a key advantage. A controller provides a centralized source of policy and desired state, enabling standardized configuration deployment across a group of devices. It can validate policy, maintain inventory and topology context, and continuously reconcile device state with the intended configuration. This reduces configuration drift, manual variation, and operational errors, while supporting repeatable changes and compliance requirements. Cisco describes software-defined networking as separating control from forwarding and enabling centralized, programmable policy and automation; these architectural characteristics underpin both device abstraction and consistent deployment. See Cisco's [Software-Defined Networking overview](#) and [Cisco DevNet](#) resources.

QUESTION NO: 65

An enterprise solution team is performing an analysis of multilayer architecture and multicontroller SDN solutions for multisite deployments. The analysis focuses on the ability to run tasks on any controller via a standardized interface. Which requirement addresses this ability on a multicontroller platform?

- A. Deploy a root controller to gather a complete network-level view.
- B. Use the East-West API to facilitate replication between controllers within a cluster.
- C. Build direct physical connectivity between different controllers.
- D. Use OpenFlow to implement and adapt new protocols.

ANSWER: B

Explanation:

Use the East-West API to facilitate replication between controllers within a cluster. A multicontroller SDN platform must operate as a coordinated distributed system rather than as a collection of independent controller instances. Controller-to-controller communication over an East-West interface enables the members of a cluster to synchronize the information needed to process requests consistently, including policy, intent, topology, endpoint, and operational state. With replicated and coordinated state, a task submitted through the platform's standardized northbound interface can be accepted and handled by an available controller without creating a different result based on which cluster member receives the request. This capability supports horizontal scaling, high availability, and continuity when a controller is unavailable or when client requests are distributed across multiple controllers. The East-West mechanism also provides the coordination foundation needed to maintain a coherent control-plane view across the deployment, which is particularly important when the architecture spans multiple sites. Cisco describes SDN as separating and centrally managing network control and policy functions, while clustered controller designs extend that management model with distributed resilience and scale. See Cisco's [Software-Defined Networking overview](#) and the [ONF SDN Architecture](#).

QUESTION NO: 66

In a redundant hub-and-spoke design with inter-spoke links, load oscillation and routing instability occur due to overload conditions. Which two design changes improve resiliency? (Choose two)

- A. Increase the number of redundant paths considered during the routing convergence calculation
- B. Eliminate links between every spoke
- C. Increase routing protocol convergence timers
- D. Increase unequal-cost parallel paths
- E. Use two links to each remote site instead of one

ANSWER: D E

Explanation:

Increasing unequal-cost parallel paths improves resiliency by allowing traffic to use multiple feasible paths rather than concentrating demand on a single lowest-metric path. In protocols and implementations that support unequal-cost multipath, such as EIGRP with the variance mechanism, traffic can be shared across qualified paths in proportion to their metrics. This creates additional forwarding capacity during periods of high utilization and makes the design less sensitive to congestion on one path. Cisco describes the EIGRP unequal-cost load-balancing behavior and feasibility requirements in its [EIGRP unequal-cost load balancing documentation](#).

Using two links to each remote site instead of one adds physical and logical path diversity between the hub and every spoke. The additional circuit supplies more available bandwidth when both links are active and retains connectivity if one circuit, interface, or provider path fails. With appropriate routing metrics and load-sharing policy, the parallel links can distribute demand predictably, reducing the likelihood that an overloaded single attachment triggers repeated route changes. This

aligns with Cisco high-availability design guidance to use redundant paths and eliminate single points of failure, as described in the [Cisco High Availability Campus Network Design Guide](#).