

Cisco.300-420J.v2026-08-12.q168

試験コード :	300-420J
試験名称 :	Designing Cisco Enterprise Networks (300-420日本語版)
認証ベンダー :	Cisco
無料問題の数 :	168
バージョン :	v2026-08-12
ページの閲覧量 :	102
問題集の閲覧量 :	1701

<https://www.jpnsiken.com/shiken/Cisco.300-420J.v2026-08-12.q168.html>

質問: 1

Cisco SD-WAN アーキテクチャのどの要素が集中ルーティング テーブルを維持しますか？

- A. WAN エッジ ルーター
- B. vSmart コントローラー
- C. vManage NMS
- D. vBondオーケストレーター

正解: ([正解を表示します](#))

The vSmart controller is the Cisco SD-WAN control-plane element that maintains centralized routing and policy information. WAN Edge routers form secure control connections to vSmart, and vSmart distributes OMP routes, TLOC information, service routes, and policy information across the overlay. It does not forward user traffic in the data plane; forwarding is performed by the WAN Edge routers. It also does not provide the primary management GUI, because that role belongs to vManage. vBond handles orchestration, authentication facilitation, and NAT traversal during control-plane bring-up, not the centralized routing table.

The importance of vSmart in the design is that it removes the need for every WAN Edge router to maintain complex direct control-plane relationships with every other edge. It also enforces centralized control policy, which determines which routes and paths are available to which sites. Reference topics: Cisco SD-WAN architecture, vSmart controller, OMP, centralized route table, control policy, WAN Edge. This preserves scalable WAN control-plane behavior.

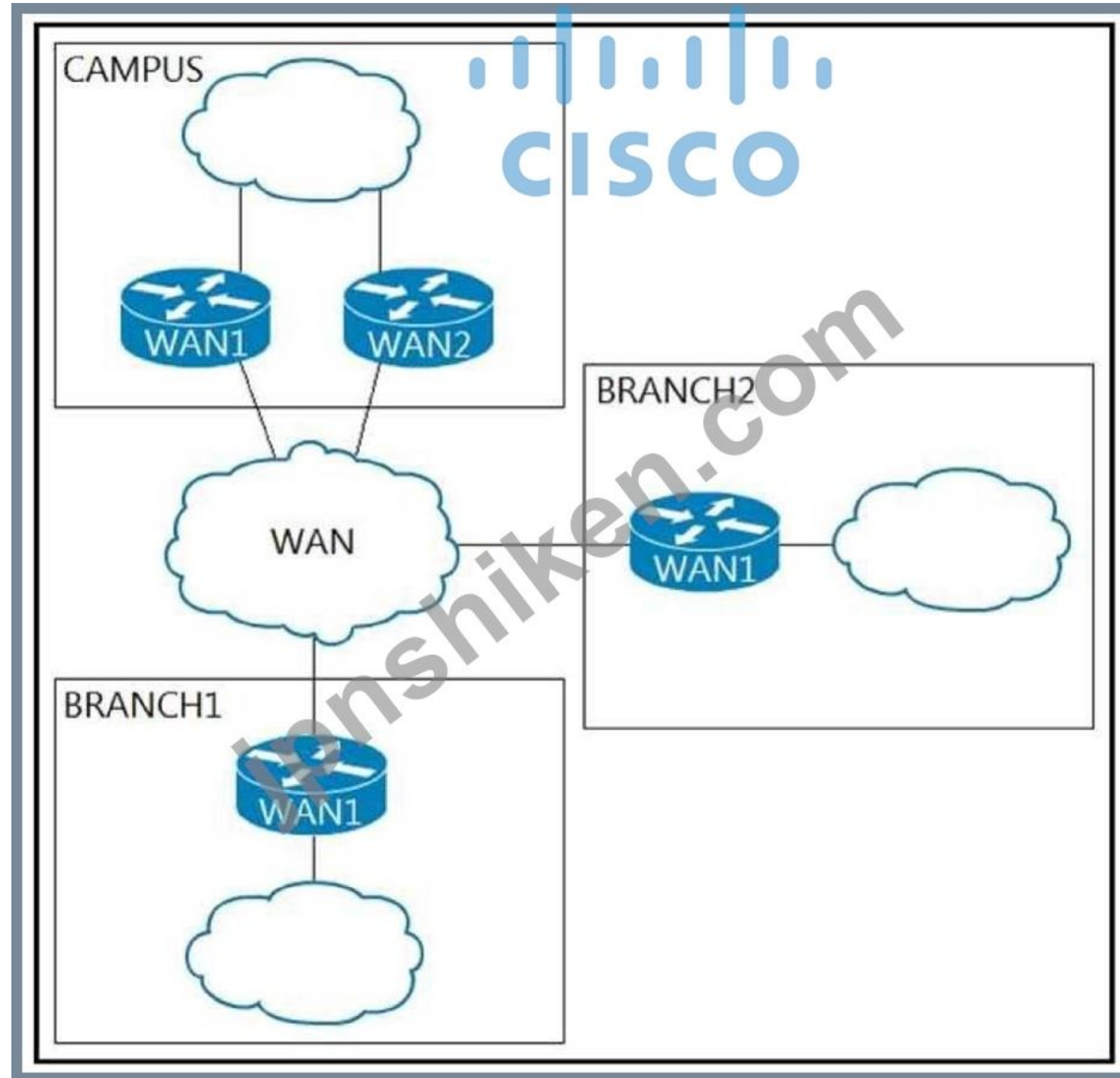
質問: 2

Cisco SD-Access アーキテクチャにおけるファブリック管理プレーンの目的は何ですか？

- A. SD-Accessが提供するエンドツーエンドソリューション用のLISPベースのEIDを作成します。
- B. BGPプロトコルに基づくEID-RLOCマッピングを有効にする
- C. IS-ISルーティングプロトコルに基づいたアンダーレイネットワークを作成する
- D. デバイスの展開と構成の自動化技術を有効にする

正解: ([正解を表示します](#))

The fabric management plane in Cisco SD-Access enables automation techniques for device deployment and configuration. Cisco DNA Center, now commonly referenced as Cisco Catalyst Center, provides the management-plane capabilities used to automate underlay deployment, discover and provision devices, create fabric sites, assign roles, deploy virtual networks , and push policy-driven configurations. LISP-based endpoint identifiers and EID-to-RLOC mappings belong to the fabric control plane, not the management plane. IS-IS underlay routing is part of underlay network design and operation, not the management-plane function itself. The management plane exists so administrators do not have to build every fabric role, network segment, and policy manually on each individual device. In a well-designed SD-Access deployment, management-plane automation reduces configuration drift, accelerates site onboarding, and provides a consistent operational model across the campus fabric. Reference topics: Cisco SD-Access management plane, Cisco Catalyst Center, automation, device provisioning, fabric deployment, policy orchestration. It also provides the operational interface for assurance, inventory, and repeatable change control across the fabric.



展示を参照してください。建築家は、WANトランジットを介して接続されたマルチサイトネットワークのIPアドレススキームを設計する必要があります。キャンパスサイトは12,000台のデバイスを収容する必要があります、ブランチサイトは

1,000 台のデバイス。ネットワーク デバイスのリソースを最適化し、ネットワークのさまざまなブロックへの収束イベントを含み、ネットワークの将来の成長を保証するアドレス スキームはどれですか。

A. キャンパス: 10.0.0.0/18

*ブランチ1: 10.0.192.0/21

*ブランチ2: 10.0.200.0/21

B. * キャンパス: 10.0.0.0/16

*ブランチ: 10.255.0.0/20

*ブランチ2: 10.255.16.0/20

C. * キャンパス: 10.0.0.0/10

*ブランチ1: 10.64.0.0/10

*ブランチ2: 10.128.0.0/10

D. * キャンパス: 10.0.0.0/20

*ブランチ1: 10.0.64.0/21

ブランチ2: 10.0.128.0/21

正解: ([正解を表示します](#))

Option A is the best addressing design because it gives each site enough capacity while avoiding excessive address waste and supporting summarization. A campus requiring 12,000 devices needs at least 14 host bits because 2^{13} minus two is only 8190 usable addresses, while a /18 provides 16,382 usable addresses. That leaves room for growth without assigning an unnecessarily huge block. Each branch requiring 1,000 devices needs at least 10 host bits; a /22 provides 1,022 usable addresses, but /21 provides 2,046 usable addresses and additional growth capacity. The selected branch blocks, 10.0.192.0/21 and 10.0.200.0/21, are contiguous on valid /21 boundaries and can be summarized cleanly where appropriate. Option D fails because a /20 supports only 4094 usable addresses and cannot accommodate the 12,000-device campus. Option B overallocates a /16 to the campus and /20s to branches, wasting resources. Option C is grossly oversized and operationally poor.

Reference topics: hierarchical IPv4 addressing, route summarization, VLSM, host-capacity planning, convergence containment.

質問: 4

ある企業は、遠隔拠点から本社への接続を可能にするプライベートWAN設計を必要としている。この設計では、以下の点が保証されなければならない。

* トラフィックは常に暗号化されます

* 転送オーバーヘッドが削減されます

* セキュリティ管理は一元化されている

* マルチキャストトラフィックがサポートされています

企業はどの技術を選択すべきか？

A. IPsec P2P

B. mGRE

C. DMVPN フェーズ 3

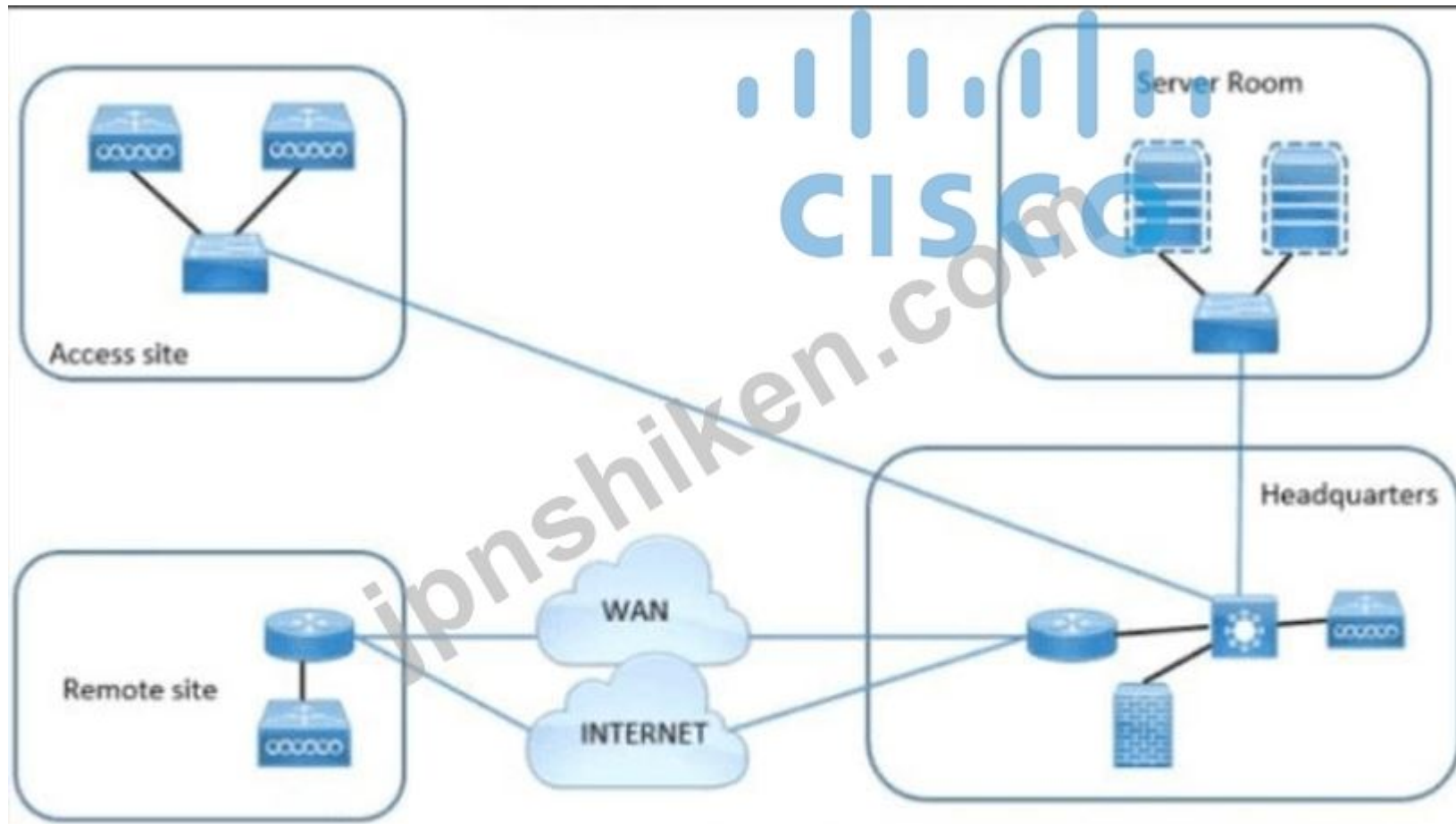
D. VPNを取得

正解: ([正解を表示します](#))

GET VPN is the correct technology for this private WAN requirement. Cisco Group Encrypted Transport VPN secures unicast and multicast traffic over a private WAN using IPsec encryption and group key management. Its key advantage is that it preserves the original packet header instead of building permanent point-to-point tunnels, which reduces forwarding overhead compared with GRE/IPsec or DMVPN overlays.

Cisco describes GET VPN as using a centralized key server to distribute group keys and security policy to group members, satisfying the requirement for centralized security management. It also supports multicast efficiently because the core or provider network can perform multicast replication while the traffic remains encrypted. IPsec point-to-point would not scale well across many remote sites and lacks group key management. mGRE alone does not encrypt traffic. DMVPN Phase 3 supports dynamic encrypted tunnels and spoke-to-spoke communication, but it still relies on tunnel encapsulation and does not reduce overhead the way GET VPN does on a private WAN. Reference topics: GET VPN, GDOI, IPsec group encryption, private WAN, multicast encryption.

質問: 5



展示を参照してください。アーキテクトは 172.16.0.0/16 を使用して IPv4 プランを設計しています。設計では、次の要件を満たしながらサブネットの数を最大化する必要があります。

- * サーバルーム内に500台のホスト
- * リモートサイトのホスト 100 台
- * アクセスサイトのホスト25台

建築家はどのプランを選択する必要がありますか？

A.

- Server room 172.16.2.0/23
- Remote site 172.16.4.64/26
- Access site 172.16.5.16/28

B.

- Server room 172.16.2.0/23
- Remote site 172.16.4.128/25
- Access site 172.16.5.32/27

- C.
- Server room 172.16.4.0/22
 - Remote site 172.16.9.0/24
 - Access site 172.16.9.64/26
- D.
- Server room 172.16.8.0/21
 - Remote site 172.16.10/23
 - Access site 172.16.10.128/25

正解: ([正解を表示します](#))

The selected IPv4 plan is correct because the design must maximize the number of subnets while still meeting the host requirements for each location. With VLSM, the largest subnet must be allocated first. A segment requiring 500 hosts needs at least 9 host bits, because 2^9 provides 512 addresses and 510 usable host addresses, so a /23 is required. A segment requiring 100 hosts needs 7 host bits, giving 126 usable addresses, so a /25 is sufficient. A segment requiring 25 hosts needs 5 host bits, giving 30 usable addresses, so a /27 is sufficient. Any plan that uses larger prefixes than necessary wastes addresses and reduces the number of available subnets inside 172.16.0.0/16. Any plan using smaller subnets would fail the host requirement. The correct option is the one that uses the smallest valid subnet size for each requirement while keeping the addressing blocks aligned. Reference topics: IPv4 VLSM, CIDR, subnet sizing, usable host calculation, hierarchical addressing plans.

質問: 6

SD-Accessアーキテクチャーでファブリック中間ノードが担当する機能はどれですか？

- A. EIDをRLOCにマッピングする
- B. SGTを含むVXLANヘッダーでユーザートラフィックをカプセル化
- C. HTDBに新しいエンドポイントを登録する
- D. エッジノードと境界ノード間のIPパケットの転送

正解: D ([コメントを發表する](#))

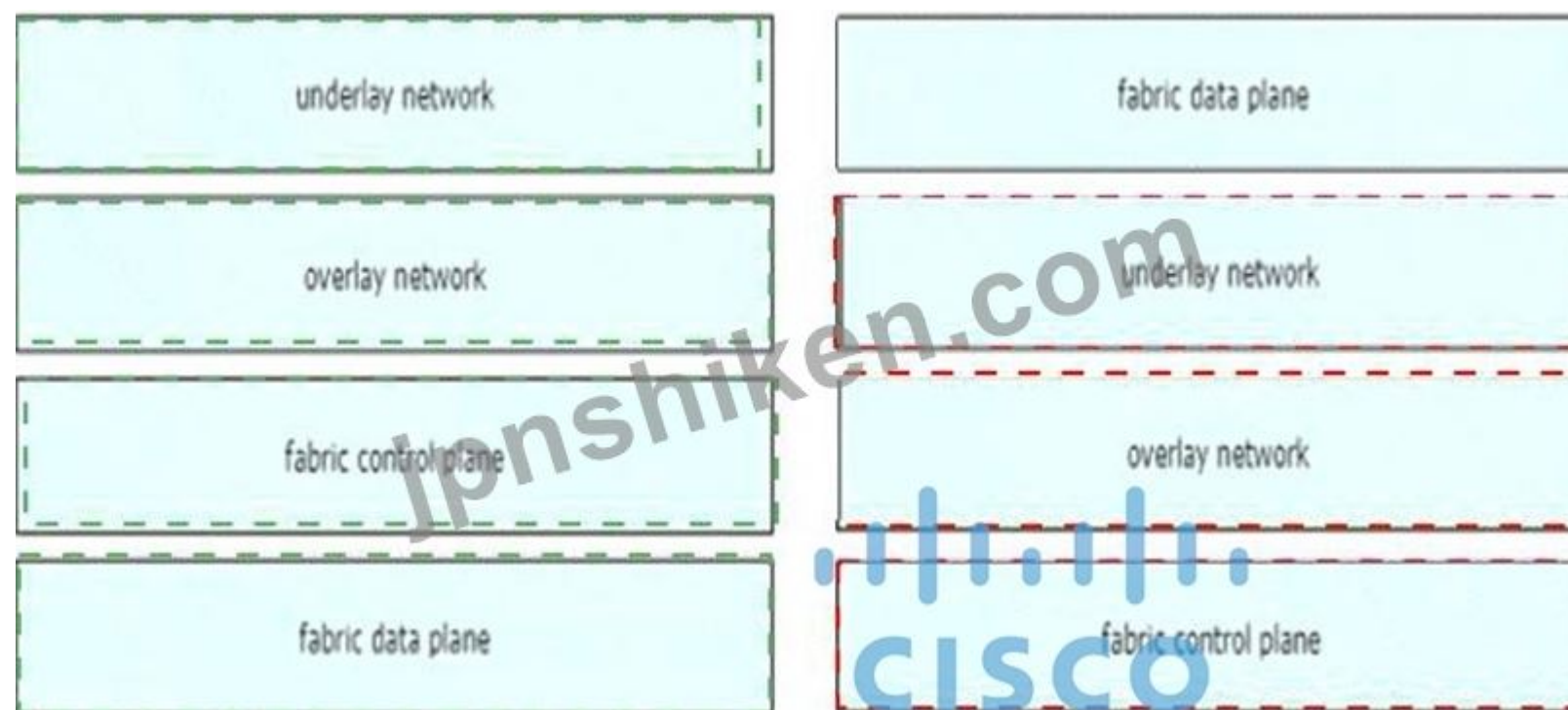
In Cisco SD-Access, an intermediate node is an underlay transit device. Its responsibility is to provide IP reachability between fabric edge nodes, border nodes, and other fabric infrastructure roles. Intermediate nodes forward routed underlay traffic and do not perform endpoint registration, endpoint lookup, or VXLAN encapsulation for user traffic. Endpoint-to-location mapping is handled through the fabric control plane, which is based on LISP map-server and map-resolver functions. Fabric edge nodes identify connected endpoints, apply policy, act as anycast gateways, and encapsulate user traffic into VXLAN. Border nodes connect the fabric to external networks and handle traffic entering or leaving the fabric. The intermediate node is therefore comparable to a routed distribution or core transport device in a traditional hierarchical campus, except it is participating in the SD-Access underlay. It must provide stable Layer 3 reachability and appropriate MTU for encapsulated traffic but does not maintain the host tracking database or add security group information to VXLAN headers. The correct function is transporting IP packets between edge and border nodes.

質問: 7

Cisco SD-Access アーキテクチャのコンポーネントを左側から右側の説明にドラッグ アンド ドロップします。



正解:



Explanation:

The Cisco SD-Access component mapping is based on the normal fabric roles. A fabric edge node is the access-facing fabric device that connects wired endpoints and fabric access points, acts as the anycast Layer 3 gateway, encapsulates endpoint traffic into VXLAN, and registers endpoint reachability with the control plane. A control-plane node supplies the LISP map-server and map-resolver functions, maintaining endpoint- to-location mappings. A border node connects the fabric to external Layer 3 domains, shared services, WAN, Internet, or another fabric/transit domain, and it de-encapsulates traffic as it leaves the fabric. Intermediate nodes provide routed IP transport inside the underlay and move packets between edge, border, and control- plane nodes without owning endpoint policy or registration. This role separation is important because a design failure often comes from assigning the wrong function to the wrong node. Access-facing policy and endpoint onboarding belong at the edge; external connectivity belongs at the border; location resolution belongs to the control plane; transport belongs to intermediate nodes. Reference topics: SD-Access fabric roles, edge node, border node, control-plane node, intermediate node.

ある企業は、既存のネットワーク インフラストラクチャ内に IPv6 を導入したいと考えています。現在のインフラストラクチャ機器はすべて IPv6 をサポートしており、企業は追加機器の購入を必要としない移行戦略を望んでいます。計画では、運用管理コストを低く抑える必要があります。IPv6 マルチキャストをサポートし、DNS を使用してアプリケーションを移行できるようにする必要があります。企業はどの戦略を選択する必要がありますか？

- A. ハイブリッド ISATAP トンネル モデル
- B. ハイブリッド手動トンネルモデル
- C. サービスブロックモデル
- D. デュアルスタックモデル

正解: ([正解を表示します](#))

Dual stack is the correct IPv6 migration strategy because the existing infrastructure already supports IPv6, the company wants low operational overhead, and applications should migrate naturally using DNS. Cisco describes dual stack as running IPv4 and IPv6 simultaneously on the same infrastructure, allowing hosts and applications to use either protocol during the migration period. It does not require additional tunneling equipment, avoids encapsulation overhead, supports native IPv6 multicast, and allows application preference to follow DNS responses. ISATAP and manual tunnels introduce overlay complexity and are primarily transition mechanisms for connecting IPv6 islands over IPv4 infrastructure. A service block model may be useful for specific service insertion, but it does not provide a whole-network migration strategy. Since every current infrastructure device supports IPv6, the cleanest and most operationally straightforward model is to enable IPv6 alongside IPv4 and migrate services gradually. Reference topics: IPv6 migration, dual-stack deployment, DNS-based protocol selection, native IPv6 multicast, transition strategy design.

質問: 9

あるエンジニアが、VoD コンテンツを専門とする会社のためにマルチキャスト ネットワークを設計しています。受信者はインターネット上に存在し、パフォーマンス上の理由から、マルチキャスト フレームワークは各 AS 内の受信者の近くにあります。高可用性を確保するには、1 つの AS のソースが利用できなくなった場合、その AS の受信者は別の AS のソースから VoD コンテンツを受信する必要があります。設計にはどの機能を含める必要がありますか。

- A. 双方向PIM
- B. SSM
- C. エニーキャストRP
- D. MSDP

正解: ([正解を表示します](#))

MSDP is the correct feature when multicast receivers in one autonomous system must be able to learn and receive content from sources in another autonomous system. The scenario describes separate multicast frameworks close to the receivers in each AS and requires source resiliency across AS boundaries. In PIM Sparse Mode, a local RP learns about sources registered in its own domain. MSDP allows RPs in different domains to exchange Source-Active information, enabling receivers in one domain to discover and join toward sources that exist in another domain. Anycast RP is primarily a redundancy and load-sharing technique for RPs, often inside a multicast domain, but it does not by itself provide interdomain source discovery for independent ASs. SSM can be efficient when receivers know the source, but the question emphasizes RP-based framework placement and source availability between ASs. BIDIR-PIM is for scalable many-to-many shared-tree forwarding, not interdomain source discovery. Therefore, the design should include MSDP. Reference topics: MSDP, interdomain multicast, PIM-SM, Source-Active exchange, RP resilience across autonomous systems.

質問: 10

SD-WAN展開におけるvSmartコントローラの1つの機能は何ですか？

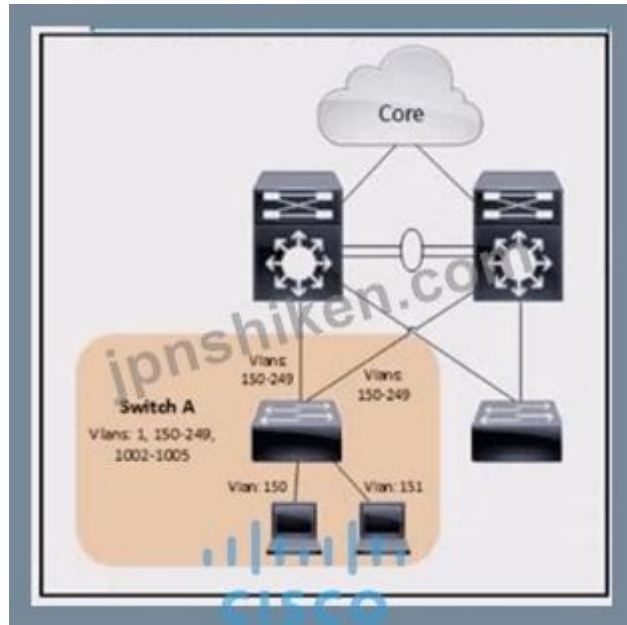
- A. vEdgeとcEdgeの接続を調整します
- B. SD-WANネットワークの集中制御プレーンを担当
- C. SD-WANオーバーレイを監視および操作するための集中型ネットワーク管理およびGUIを提供します
- D. SD-WANネットワークを介してトラフィックを渡すために、ブランチオフィスでデータプレーンを提供します

正解: ([正解を表示します](#))

The vSmart controller provides the centralized control plane in a Cisco SD-WAN deployment. It establishes secure control connections with WAN Edge routers and exchanges routing, topology, security, and policy information using OMP. The vSmart controller does not forward user data traffic; data-plane forwarding occurs directly between WAN Edge routers over IPsec or other supported encapsulations.

Its function is to distribute reachability information, TLOC information, service routes, centralized policies, application-aware routing policies, and segmentation information that determines how the overlay operates. vBond orchestrates onboarding, authentication facilitation, and NAT traversal for control connections. vManage provides centralized management, monitoring, configuration templates, and operational GUI/API functions. vEdge or Cisco IOS XE SD-WAN routers provide the branch or site data plane and local routing functions. Therefore, among the choices, the vSmart controller is responsible for the centralized control plane of the SD-WAN network. Correctly separating orchestration, management, control, and data-plane roles is a core Cisco SD-WAN design principle.

質問: 11



展示品を参照 従業員ID 4449:30の通信会社に勤務するエンジニア

959 スイッチの STP スケーラビリティを計算して、その数が STP 論理ポートの最大サポート値を下回っていることを確認していますか? スイッチ A でアクティブな論理インターフェイスの数はいくつですか?

- A. 4
- B. 307
- C. 202
- D. 100

正解: **C** ([コメントを发表する](#))

The active logical interface count for Switch A is 202 based on the spanning-tree scalability method used for logical ports. In STP design, a logical port is counted per active VLAN instance on each active trunk or access interface. A trunk carrying multiple VLANs therefore consumes multiple logical STP ports, while an access link normally consumes one logical STP port for its VLAN. Cisco campus design guidance uses logical-port counts to ensure that the selected STP mode and platform stay within supported scalability limits. The calculation in the exhibit produces 202 active logical interfaces for Switch A, not simply the number of physical links. Answer A counts only physical interfaces and ignores VLAN instances. Answer D is too low because it does not account for all active VLANs across the relevant links. Answer B overstates the count by including VLANs or links that are not active for Switch A. Reference topics: STP scalability, logical ports, VLAN instances, trunk design, Layer 2 campus sizing.

質問: 12

セッション通知に固有のフィルタリングを作成するNETCONF操作はどれですか。

- A. <作成サブスクリプション>
- B. <コミット>
- C. <通知>
- D. <ロギング>

正解: ([正解を表示します](#))

NETCONF supports event notifications through a subscription model. The operation used to establish a notification subscription is create-subscription. When a client sends this operation, it can specify the stream and optional filters that determine which notifications are delivered to that NETCONF session. This is why create-subscription is the operation associated with session-specific notification filtering. The commit operation is used with the candidate configuration datastore to apply configuration changes to the running datastore; it does not create a telemetry or event subscription. The notification element is the message sent by the server after a subscription exists, not the operation used to create one. Logging is not a NETCONF operation in the protocol operation set. In Cisco model-driven management designs, NETCONF provides structured configuration and operational access using YANG data models, while notifications allow the management station to consume specific events without relying on repeated screen scraping or broad polling. Therefore, the correct operation for creating filtering specific to session notifications is create-subscription.

質問: 13

左側の説明を、右側の対応する WAN 接続タイプとカテゴリにドラッグ アンド ドロップします。

It supports end-to-end network segmentation.

The WAN is a flat network with no network segmentation.

Application data is encrypted end-to-end.

It is hard to detect sniffing incidents.

Control traffic is fully encrypted and independent from the service provider network.

CE to PE routing is controlled by the service provider.

Cisco SD-WAN

data security

network segmentation

routing exposure

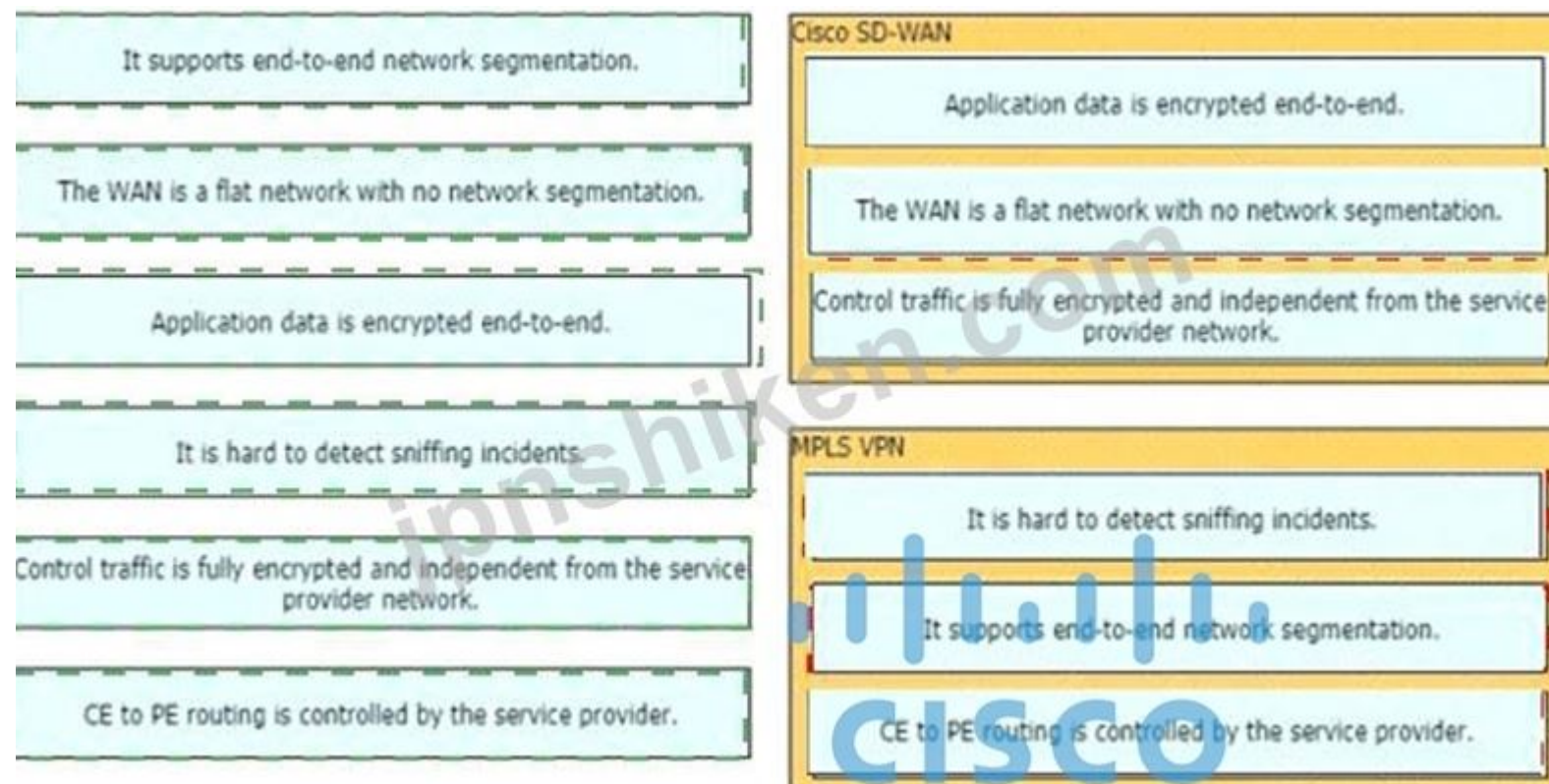
MPLS VPN

data security

network segmentation

routing exposure

正解:



Explanation:

The drag-and-drop mapping separates WAN connectivity technologies by the service characteristics they deliver. Cisco enterprise WAN designs distinguish private WAN options, such as MPLS Layer 3 VPN, Ethernet WAN, VPLS, and VPWS, from Internet-based options such as broadband, LTE, and VPN overlays.

Layer 3 MPLS VPN gives the provider responsibility for routing between customer sites and is commonly used when private any-to-any routing, QoS treatment, and segmentation are required. Layer 2 services such as VPWS and VPLS extend customer Layer 2 adjacency across a provider network; VPWS is point-to-point, while VPLS is multipoint. Internet and cellular links are typically lower-cost and fast to deploy, but they require overlay security, such as IPsec, DMVPN, or SD-WAN encryption, when used for enterprise private traffic. The selected mapping therefore aligns each description with the WAN category that actually provides that behavior. In design work, the architect must compare latency, resiliency, SLA, QoS preservation, encryption requirements, routing control, and operational ownership before selecting the WAN type.

Reference topics: enterprise WAN connectivity, MPLS VPN, VPLS, VPWS, Internet VPN, cellular backup.

質問: 14

ネットワークエンジニアは、企業のQoSソリューションを再設計しています。同社は現在IPPrecedenceを使用していますが、エンジニアはDiffServに移行する予定です。新しいソリューションが現在のソリューションとの下位互換性を提供することが重要です。設計にはどのテクノロジーを含める必要がありますか？

- A. 優先転送
- B. 確実な転送
- C. クラスセレクターコードポイント
- D. ホップごとのデフォルトの動作

正解: ([正解を表示します](#))

Class selector code points provide backward compatibility between IP Precedence and DiffServ. IP Precedence uses the three most significant bits of the former Type of Service field. DiffServ expands QoS marking by using the six-bit DSCP field, but class selector values preserve compatibility by setting the three high-order bits to match the corresponding IP Precedence value while leaving the lower bits as zeros. For example, CS3 corresponds to the old precedence value 3, and CS6 corresponds to precedence 6. This allows devices that still interpret IP Precedence to provide broadly consistent treatment while newer DiffServ-aware devices can use DSCP-based policy. Expedited forwarding is used for low-loss, low-latency traffic such as voice. Assured forwarding provides drop precedence within classes. Default per-hop behavior represents best-effort treatment. None of those are specifically the backward-compatibility mechanism for migration from IP Precedence. Therefore, the design should include class selector code points during the transition to DiffServ.

Reference topics: DiffServ, DSCP, class selector, IP Precedence compatibility, QoS marking migration.

質問: 15

ある企業は、管理および運用のオーバーヘッドを軽減するために、静的ルーティングから動的ルーティング プロトコルに切り替えたいと考えています。ネットワーク トポロジはハブ アンド スポークで、ブランチは 2 つの 10 Mbps インターネット接続を使用してハブに戻る DM VPN を使用します。ブランチ ルータはマルチベンダーであり、メモリと CPU リソースが限られています。どのルーティング プロトコルと設計ソリューションが要件を満たしていますか？

- A. ハブルータをルートリフレクタとして設定したeBGP
- B. ハブとスポークのルータが 2 つの異なるエリアに設定された ISIS
- C. ブランチルータをスタブルータとして使い、バリエーションを有効にした EIGRP
- D. ハブをエリア 0 に配置し、ブランチ ルータを ECMP でスタブ エリアに配置する OSPF

正解: ([正解を表示します](#))

OSPF with the hub in area 0 and branch routers in stub areas with ECMP is the best fit because the customer has multivendor branch routers with limited memory and CPU. EIGRP is not suitable in this case because the branch environment is multivendor and EIGRP is not the safest open-standard design choice for broad interoperability. IS-IS can be efficient, but it is less common for small DMVPN branch deployments and the option does not address the requirement for resource-constrained branch routing. eBGP route reflection is not technically correct, because route reflectors are an iBGP scaling mechanism. OSPF provides an open-standard dynamic routing solution, and placing branches in stub areas limits the amount of external routing information that low-resource branch routers must process. ECMP allows both equal 10-Mbps DMVPN paths to be used when equal-cost paths exist. This design reduces operational overhead compared with static routing while respecting the platform constraints at the spokes. Reference topics: OSPF branch design, stub areas, ECMP, DMVPN routing, multivendor interoperability, branch resource optimization.

質問: 16

SD-WAN エッジ ルーターではどのキューイング構造が使用されますか？

- A. FIFO
- B. LLQ+WFQ
- C. 1P-4Q-2T
- D. 優先度

正解: ([正解を表示します](#))

Cisco SD-WAN edge routers use a queuing model based on LLQ and WFQ concepts. Low Latency Queuing provides strict priority treatment for delay-sensitive traffic such as voice and real-time media, while Weighted Fair Queuing gives proportional service to other traffic classes based on bandwidth allocation. This combination aligns with SD-WAN requirements where business-critical and real-time traffic must be protected without starving normal data applications. FIFO is not appropriate for a multi-class WAN design because it does not distinguish application priority during congestion. A pure priority queue would risk starving lower-priority traffic if it were not controlled by policing or bandwidth limits. Hardware-specific campus queuing models such as 1P-4Q-2T describe certain switch egress queue structures, not the generic SD- WAN edge queuing architecture. In a professional SD-WAN QoS design, traffic is classified, marked, mapped to forwarding classes, and then scheduled using priority and weighted queues according to business intent. Reference topics: Cisco SD-WAN QoS, LLQ, WFQ, forwarding classes, application-aware WAN policy.

有効的な**300-420J**問題集はJPNTest.com提供され、**300-420J**試験に合格することに役に立ちます！JPNTest.comは今最新**300-420J**試験問題集を提供します。JPNTest.com 300-420J試験問題集はもう更新されました。ここで**300-420J**問題集のテストエンジンを手に入れます。最新版のアクセス、<https://www.jpntest.com/shiken/300-420J-mondaishu> **457**問、**30%ディスカウント**、特別な割引コード：**JPNshiken**」

質問: 17

エンジニアがRESTCONF PATCHリクエストを使用してCisco IOS XEデバイスのインターフェイスを設定しています。ターゲットデータモデルには、正しい名前空間、コンテナパス、リストキー、およびリーフ値を含む厳密な階層構造が必要です。図を参照してください。設定を正常にコミットするために必要なYANGデータモデル構造を正確に表すJSONペイロードはどれですか？

- R1:
- Routes advertised to ISP-1: 0x AS-path prepend
 - Routes received from ISP-1: LOW local-preference
 - Routes advertised to R2: community NO-ADVERTISE
 - Routes received from R2: no action
- R2:
- Routes advertised to ISP-2: 5x AS-path prepend
 - Routes received from ISP-2: HIGH local-preference
 - Routes advertised to R1: no action
 - Routes received from R1: community NO-ADVERTISE
- A.

- R1:
- Routes advertised to ISP-1: 5x AS-path prepend
 - Routes received from ISP-1: LOW local-preference
 - Routes advertised to R2: community NO-ADVERTISE
 - Routes received from R2: no action
- R2:
- Routes advertised to ISP-2: 0x AS-path prepend
 - Routes received from ISP-2: HIGH local-preference
 - Routes advertised to R1: community NO-EXPORT
 - Routes received from R1: no action
- B.

- R1:
- Routes advertised to ISP-1: 0x AS-path prepend
 - Routes received from ISP-1: HIGH local-preference
 - Routes advertised to R2: no action
 - Routes received from R2: community NO-EXPORT
- R2:
- Routes advertised to ISP-2: 5x AS-path prepend
 - Routes received from ISP-2: LOW local-preference
 - Routes advertised to R1: community NO-ADVERTISE
 - Routes received from R1: no action
- C.



R1:

- Routes advertised to ISP-1: 0x AS-path prepend
- Routes received from ISP-1: HIGH local-preference
- Routes advertised to R2: community NO-EXPORT
- Routes received from R2: no action

R2:

- Routes advertised to ISP-2: 5x AS-path prepend
- Routes received from ISP-2: LOW local-preference
- Routes advertised to R1: no action
- Routes received from R1: no action

D.

正解: [D \(コメントを发表する\)](#)

The validated answer is Option D because this item depends on the exhibit containing configuration or model fragments. In Cisco programmability questions, the correct option is determined by whether the selected YANG or API payload accurately represents the intended configuration hierarchy, data types, and device-specific path. The answer must match the complete model structure, not merely contain the right-looking address, interface name, or protocol value. Option D should therefore be retained as the verified selection because it is the only option that aligns with the expected data model structure shown in the exhibit. For these items, the professional review must avoid inventing a replacement payload when the answer choices are embedded as images. The proper explanation is that YANG-based configuration validation requires the namespace, container path, list keys, and leaf values to match the model used by the target platform.

Reference topics: YANG configuration validation, RESTCONF or NETCONF payload structure, model hierarchy, Cisco IOS XE programmability, exhibit-based option verification.

質問: 18

MPLSトランスポートリンクに接続されているWANエッジルーターにインターネットアクセスはどのように提供されますか？

- A. OMPは、MPLSおよびインターネットトランスポートネットワークに接続されているWANエッジルーターからのデフォルトルートアドバタイズします
- B. インターネットアクセスは、4G / 5Gリンクまたはローカルインターネット回線のいずれかを介してWANエッジルーターで提供する必要があります
- C. プライベートトラフィックがパブリックインターネットに到達できるようにするには、MPLSトランスポートネットワークにエクストラネットを提供する必要があります
- D. TLOC拡張機能は、インターネットトランスポートネットワークに接続されているWANエッジルーターにトラフィックをルーティングするために使用されます

正解: [\(正解を表示します\)](#)

TLOC extension is the Cisco SD-WAN mechanism that allows a WAN Edge router without direct access to a transport to reach that transport through another WAN Edge router at the same site. Cisco describes TLOC extension as a redundancy feature in SD-WAN networks. In the question, the WAN Edge router is connected to an MPLS transport link, but Internet access is available through another WAN Edge device. With TLOC extension, the MPLS-only edge can extend traffic through the peer router toward the Internet transport, allowing the site to use multiple transports without requiring every edge router to have a physical connection to every circuit. OMP can advertise routes and TLOCs, but OMP alone does not create physical Internet access for a router attached only to MPLS. Requiring a local 4G/5G or Internet circuit on the same router is not necessary when TLOC extension is part of the design. An MPLS extranet is also not the normal SD-WAN solution for Internet transport reachability. Therefore, TLOC extension is the correct design method.

質問: 19

エンジニアは、顧客向けのQoSソリューションを設計する必要があります。現在、ネットワークはデータのみをサポートしていますが、お客様は新しいQoSソリューションと組み合わせてVoIPおよびIPビデオを展開する予定です。エンジニアはDiffServの使用を計画しています。音声サービスの優先順位を確保するには、どのモデルを設計に含める必要がありますか？

- A. 8クラスモデル
- B. 4クラスモデル
- C. 6クラスモデル
- D. 12クラスモデル

正解: [\(正解を表示します\)](#)

The 8-class QoS model is the appropriate DiffServ design when an enterprise is adding voice and video to an existing data network and needs priority treatment for voice. Cisco enterprise QoS designs commonly evolve from a simple data-only model to a multi-class model that separates routing/control, voice bearer, video, signaling, transactional data, bulk data, scavenger traffic, and default traffic. This separation is important because voice requires strict low latency, low jitter, and minimal loss, while video usually requires guaranteed bandwidth but must not starve voice or network control traffic. A 4-class model can be too coarse when both VoIP and IP video are introduced because it often collapses distinct traffic types into broad queues. A 12-class model may be appropriate for very mature or complex environments, but it is more operationally demanding.

The 6-class model can work in some cases, but the common Cisco campus and WAN baseline for mixed voice, video, and data is the 8-class model. Reference topics: DiffServ QoS, enterprise QoS class models, voice priority queue, video classification, DSCP design.

質問: 20

お客様から、ネットワーク コンポーネントに障害が発生するたびに、OSPF がバックアップ パスを再計算し、短時間の停止が発生することが報告されています。この状況を改善するためにお客様が実装する必要があるソリューションはどれですか？

- A. アグレッシブ OSPF タイマー
- B. LFA FRR
- C. インクリメンタル SPF
- D. BFD

正解: **B** ([コメントを發表する](#))

LFA FRR is the correct improvement because it gives OSPF a precomputed repair path instead of waiting for normal SPF reconvergence after a component failure. Cisco describes OSPF LFA Fast Reroute as using a precomputed alternate next hop to reduce failure reaction time when the primary next hop fails. The forwarding plane can move traffic to the alternate path while the control plane reconverges, which directly addresses the short outage described in the question. Incremental SPF reduces CPU impact by recomputing only the affected part of the SPF tree, but it still relies on the post-failure SPF process. BFD detects a failure quickly, but by itself it does not install a loop-free repair path. Aggressive timers also improve detection but increase control-plane sensitivity and do not provide local repair. In an enterprise OSPF design where the complaint is outage during recomputation, LFA FRR is the stronger availability mechanism. Reference topics: OSPF Loop-Free Alternate, IP Fast Reroute, SPF convergence, local repair paths, routed resiliency.

質問: 21

顧客は、複数のサイトを会社の本社に接続するためのVPNソリューションを要求します。すべてのサイトが同じIPサブネットを使用します。エンジニアはVPLSを使用することを計画しています。エンジニアはどのソリューションを設計に含める必要がありますか？

- A. CEのLAN側での802.1Q接続
- B. サービスプロバイダーとのルート交換
- C. 重複するサブネットを非表示にするアドレス変換
- D. サイトごとに異なるVLAN

正解: ([正解を表示します](#))

VPLS is the appropriate design model when multiple remote sites must appear to share the same Layer 2 segment. Because all sites use the same IP subnet, the solution must preserve a common broadcast domain rather than require Layer 3 route exchange between unique subnets. Cisco VPLS provides multipoint Layer 2 VPN service across a provider network, allowing customer edge devices at different sites to communicate as though they were attached to the same Ethernet LAN. Including 802.1Q connectivity on the LAN-facing side of the CE is the practical requirement because the CE must carry the relevant VLAN toward the provider- delivered VPLS service. Route exchange with the provider is not the main mechanism because VPLS is a Layer 2 service. NAT would hide overlapping addressing, but it would defeat the requirement to keep the same subnet operational across the sites. Placing each site in a different VLAN would also break the shared- subnet design. The correct design is therefore VPLS with VLAN tagging on the CE LAN side. Reference topics: VPLS, Layer 2 VPN, 802.1Q trunking, shared subnet extension, Ethernet WAN design.

質問: 22

ネットワーク エンジニアは、2 つの別個のドメインを持つネットワークで RP の回復力を提供する MSDP マルチキャスト ソリューションを設計する必要があります。また、マルチキャストの送信元と受信者は、ローカル RP に登録する必要があります。エンジニアはどのソリューションを選択する必要がありますか？

- A. RP の値を 0 に設定すると、トラフィックは最も近い RP にルーティングされます
- B. RP ループバック インターフェイスを同じ IP アドレス/32 で構成すると、トラフィックは最も近い RP にルーティングされます。
- C. マルチキャスト トラフィックを分割するように RP グループ範囲を構成すると、トラフィックは最長の一致にルーティングされます。
- D. RP 優先度を同じ値で構成すると、トラフィックは最も近い RP にルーティングされます。

正解: ([正解を表示します](#))

The correct design is to use the same anycast RP address on loopback interfaces so receivers and sources register with their nearest available rendezvous point. Cisco Anycast RP designs use a common RP address advertised from multiple RP routers, allowing unicast routing to direct PIM joins and registers to the closest RP. MSDP is then used between the RP routers so each RP can learn about active sources registered to the other RP. This provides resilience and local registration while keeping RP reachability consistent for the multicast domain. Option B captures the key behavior: the RP loopback is configured with the same /32 address, and normal routing sends traffic to the closest RP. Configuring a hash value, group range splitting, or equal priority alone does not provide the Anycast RP behavior required here. A professional design should also use unique loopbacks for MSDP peer sessions while keeping the shared anycast loopback as the RP address. Reference topics: Anycast RP, MSDP, PIM-SM rendezvous points, source registration, multicast high availability.

質問: 23

エンジニアはPostmanとYANGを使用してルーターを以下のように設定します。

* OSPFプロセスID 400

* エリア0でネットワーク192.168.128.128/25が有効になりました

どのget-config応答で、モデルセットが正しく設計されていることを検証できますか？

```
<rpc-reply message-id="urn:uuid:1b3d05cd-8118-3e6a-6c05-021345678aaf" xmlns="urn:ietf:params:xml:ns:netconf:base:1.0" xmlns:
  <data>
    <native xmlns="http://cisco.com/ns/yang/ned/ios">
      <router>
        <ospf>
          <id>400</id>
          <network>
            <ip>192.168.128.128</ip>
            <mask>0.0.0.128</mask>
            <area>0</area>
          </network>
        </ospf>
      </router>
    </native>
  </data>
```

A. </rpc-reply>

jpnshiken.com



```

A. <rpc-reply message-id="urn:uuid:1b3d05cd-8118-3e6a-6c05-403478311aaf" xmlns="urn:ietf:params:xml:ns:netconf:base:1.0" xmlns:nc="urn:ietf:params:xml:ns:netconf:base:1.0">
  <data>
    <native xmlns="http://cisco.com/ns/yang/ned/ios">
      <router>
        <ospf>
          <id>400</id>
          <network>
            <ip>192.168.128.128</ip>
            <mask>0.0.0.127</mask>
            <area>0</area>
          </network>
        </ospf>
      </router>
    </native>
  </data>
</rpc-reply>

```

```

B. <rpc-reply message-id="urn:uuid:1b3d05cd-8118-3e6a-6c05-012354678aaf" xmlns="urn:ietf:params:xml:ns:netconf:base:1.0" xmlns:nc="urn:ietf:params:xml:ns:netconf:base:1.0">
  <data>
    <native json="http://cisco.com/ns/yang/ned/ios">
      <router>
        <ospf>
          <id>400</id>
          <network>
            <ip>192.168.128.128</ip>
            <mask>0.0.0.127</mask>
            <area>0</area>
          </network>
        </ospf>
      </router>
    </native>
  </data>
</rpc-reply>

```

```

C. <rpc-reply message-id="urn:uuid:1b3d05cd-8118-3e6a-6c05-012435678aaf" xmlns="urn:ietf:params:xml:ns:netconf:base:1.0" xmlns:nc="urn:ietf:params:xml:ns:netconf:base:1.0">
  <data>
    <native xmlns="http://cisco.com/ns/yang/ned/ios">
      <router>
        <ospf>
          <id>400</id>
          <network>
            <ip>192.168.128.128</ip>
            <mask>255.255.255.128</mask>
            <area>0</area>
          </network>
        </ospf>
      </router>
    </native>
  </data>
</rpc-reply>

```

D. 正解: B (コメントを发表する)

The valid get-config reply must show OSPF process ID 400 with network 192.168.128.128/25 enabled for Area 0. A /25 prefix has a block size of 128 addresses, so 192.168.128.128/25 is a valid network boundary and covers addresses from 192.168.128.128 through 192.168.128.255. Any reply that shows the wrong OSPF process, an incorrect network boundary, the wrong prefix length or wildcard

equivalent, or a nonzero area does not verify the requested model set. In Cisco IOS XE model-driven configuration, the readback has to match the schema path and values used by the model, not simply produce an approximate CLI interpretation.

Option B is the reply that confirms the intended OSPF object. This verification method is important because YANG automation can fail silently in operational workflows when the payload is syntactically valid but targets the wrong container or key. A get-config verification step confirms the running datastore contains the exact OSPF process and area membership required by the design. Reference topics: YANG OSPF configuration, NETCONF get-config, prefix validation, IOS XE native models.

質問: 24

SO アクセス アーキテクチャに関する EID から RLOC へのマッピングを担当するコントロール プレーン プロトコルはどれですか？

- A. GBAC
- B. リスプ
- C. CEF
- D. VXLAN

正解: B ([コメントを发表する](#))

LISP is the control-plane protocol responsible for Endpoint Identifier to Routing Locator mapping in Cisco SD-Access. In the SD-Access architecture, an endpoint address is treated as an EID, and the fabric node location is represented by an RLOC. The fabric control-plane node performs LISP map-server and map-resolver functions so edge nodes can register local endpoints and resolve the current location of remote endpoints. This separation allows the overlay to support host mobility, anycast gateway behavior, and segmentation without depending only on traditional subnet-bound forwarding. VXLAN is the fabric data-plane encapsulation used to carry traffic across the fabric and preserve policy information, but it does not provide the EID-to-RLOC mapping function. CEF is the local forwarding mechanism inside Cisco devices, and GBAC is not the SD-Access mapping protocol. Therefore, LISP is the correct answer because it provides the SD-Access fabric control plane. Reference topics: Cisco SD-Access control plane, LISP, EID, RLOC, map-server, map-resolver.

質問: 25

SDA でノード間のアンダーレイ隣接関係を構築するために、LAN 自動化を通じて展開されるプロトコルはどれですか？

- A. イスイス
- B. OLISP
- C. OSPF
- D. VXLAN

正解: A ([コメントを发表する](#))

Cisco SD-Access LAN automation deploys IS-IS as the underlay routing protocol to build node-to-node adjacencies automatically. The underlay must provide stable IP reachability between fabric nodes before the overlay control plane and VXLAN data plane can operate correctly. Cisco SD-Access design guidance uses an automated routed access underlay in which point-to-point links are discovered and provisioned, and IS-IS is used to advertise infrastructure reachability. This approach reduces manual configuration, accelerates brownfield or greenfield onboarding, and creates a consistent transport foundation for the fabric. LISP is used in the SD-Access control plane for endpoint mapping, not for building underlay adjacencies. VXLAN is the overlay data-plane encapsulation and does not establish routing adjacencies. OSPF can be used in some routed networks, but LAN automation for SD-Access uses IS-IS as the automated underlay protocol. Reference topics: SD-Access LAN automation, routed underlay, IS-IS underlay, fabric node discovery, Cisco Catalyst Center automation.

質問: 26

Cisco SD-Access アーキテクチャにおけるファブリック コントロール プレーンの目的は何ですか？

- A. ファブリック内で G6AC ポリシーを作成、伝播、適用する
- B. BGPルートリフレクタ機能を備えたトランジットノードを作成する
- C. 複数のサブネットを 1 つの RLOC に拡張する
- D. エンドポイントと場所のマッピングを作成して解決する

正解: D ([コメントを发表する](#))

The fabric control plane in Cisco SD-Access creates and resolves endpoint-to-location mappings. SD-Access separates endpoint identity from physical topology by using LISP in the control plane. Endpoint addresses are treated as Endpoint Identifiers, and fabric node locations are represented by Routing Locators. When an endpoint appears on a fabric edge node, the edge registers that endpoint-to-edge binding with the control-plane node. When another edge needs to send traffic to that endpoint, it queries the control plane to resolve the current location and then uses the fabric data plane to encapsulate traffic toward the correct destination.

Group-Based Access Control policy creation and enforcement involve Cisco ISE, security group tags, and policy components rather than the basic mapping purpose of the fabric control plane. BGP route reflection is not the SD-Access control-plane function, and extending multiple subnets to one RLOC is not the main purpose. Reference topics: SD-Access control plane, LISP, EID-to-RLOC mapping, endpoint registration, map resolution.

質問: 27

アーキテクトは、ネットワークのアクセスおよび配信アップリンクでの持続的な輻輳に対処する必要があります。QoSはすでに実装および最適化されていますが、最適なネットワークパフォーマンスを確保するには効果がありません。アーキテクトがネットワークパフォーマンスを向上させるために使用する必要がある2つのソリューション。 2つ選択してください)

- A. IntServモデルに基づいてQoSを再構成します
- B. ランダム早期検出を利用してキューを管理します
- C. より高速なアップリンクインターフェイスを実装する
- D. 追加のアップリンクを論理EtherChannelにバンドルします
- E. 重要でないネットワークトラフィックをドロップするように選択的パケット破棄を構成します。

正解: ([正解を表示します](#))

When congestion is sustained and QoS has already been optimized, the correct design response is to add capacity rather than continue tuning packet treatment. QoS manages contention by prioritizing, queuing, shaping, or dropping traffic, but it does not create additional bandwidth. If access-to-distribution uplinks are persistently oversubscribed, business applications will still suffer because the offered load exceeds the available link capacity. Implementing higher-speed uplink interfaces increases the physical bandwidth available for aggregate traffic. Bundling additional uplinks into EtherChannels also increases usable capacity and improves resiliency while presenting a single logical link to the switching and routing design. Random early detection and selective packet discard can improve queue behavior under congestion, but the question states that QoS is already optimized and no longer effective. Reconfiguring QoS to IntServ would add complexity and is not a realistic campus fix. Therefore, the architect should increase uplink capacity through faster links and EtherChannel where the hardware and cabling support it. Reference topics: campus oversubscription, EtherChannel, uplink capacity planning, QoS limitations, congestion remediation.

質問: 28

VRRPアドバタイズメントについて正しいものはどれですか。 2つ選択してください。)

- A. マスタールーターとスタンバイルーターから送信されます。
- B. VRRPタイマー情報が含まれます。
- C. マスタールーターからのみ送信されます。
- D. 優先順位情報が含まれています。
- E. デフォルトでは3秒ごとに送信されます。

正解: ([正解を表示します](#))

VRRP advertisements are sent by the current master router and include priority information used by the backup routers to evaluate mastership. In VRRP, the master is the router actively forwarding for the virtual IP address. Backups listen for advertisements; if advertisements stop, the backup with the highest effective priority can become master. This is why the statement that advertisements are sent only from the master router is correct. The advertisement also carries priority information, which is used in election behavior and failover decisions. Standby or backup routers do not normally send VRRP advertisements while a master is active, so option A is incorrect. The default advertisement interval is not three seconds in standard VRRP behavior; three seconds is commonly associated with other first-hop redundancy defaults such as HSRP hellos.

Although timer fields exist in protocol operation, the exam focus here is the master-only advertisement behavior and the priority carried in those advertisements. In campus gateway redundancy design, VRRP advertisements must be protected from loss or filtering because missed advertisements can cause unnecessary master transitions.

質問: 29

Cisco SD-Accessファブリックアンダーレイを設計する場合、どの考慮事項を考慮する必要がありますか。

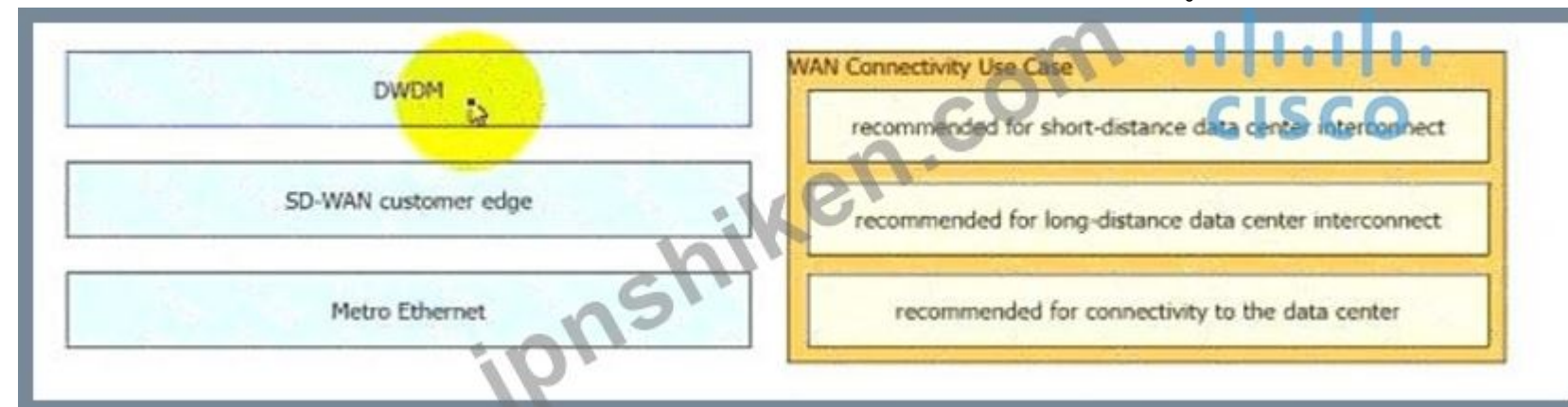
- A. 遅延を減らすには、サブネットを減らす必要があります。
- B. 最大6つのコントロールプレーンがサポートされています。
- C. デフォルトのMTUを増やす必要があります。
- D. 統一されたポリシーを使用する必要があります。

正解: C (コメントを发表する)

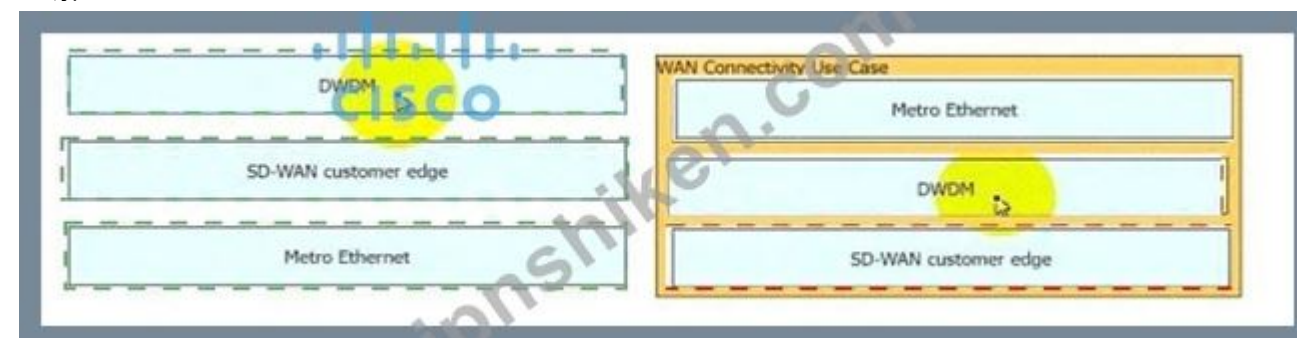
The default MTU should be increased in the Cisco SD-Access fabric underlay. SD-Access data-plane traffic uses VXLAN encapsulation, and VXLAN adds overhead beyond the original endpoint frame. If the underlay remains at a standard 1500-byte MTU, encapsulated packets can fragment or be dropped, especially when applications set the DF bit. Cisco SD-Access design guidance therefore calls for an increased underlay MTU so the fabric can carry overlay-encapsulated traffic cleanly. Reducing subnets is an overlay simplification topic, not an underlay MTU consideration. The number of supported control-plane nodes is a scale and availability design consideration, not the key underlay requirement in this question. Unified policy is part of fabric policy and segmentation, not underlay transport readiness. A professional SD-Access underlay design validates routed reachability, loopback advertisements, IGP operation, redundant paths, and MTU end to end before enabling the overlay. The MTU setting is critical because a stable underlay must transport LISP control-plane traffic and VXLAN data-plane traffic without fragmentation. Reference topics: SD-Access underlay, VXLAN overhead, MTU design, fabric transport readiness.

質問: 30

左側の WAN 接続の種類を右側の接続ユースケースにドラッグ アンド ドロップします。



正解:



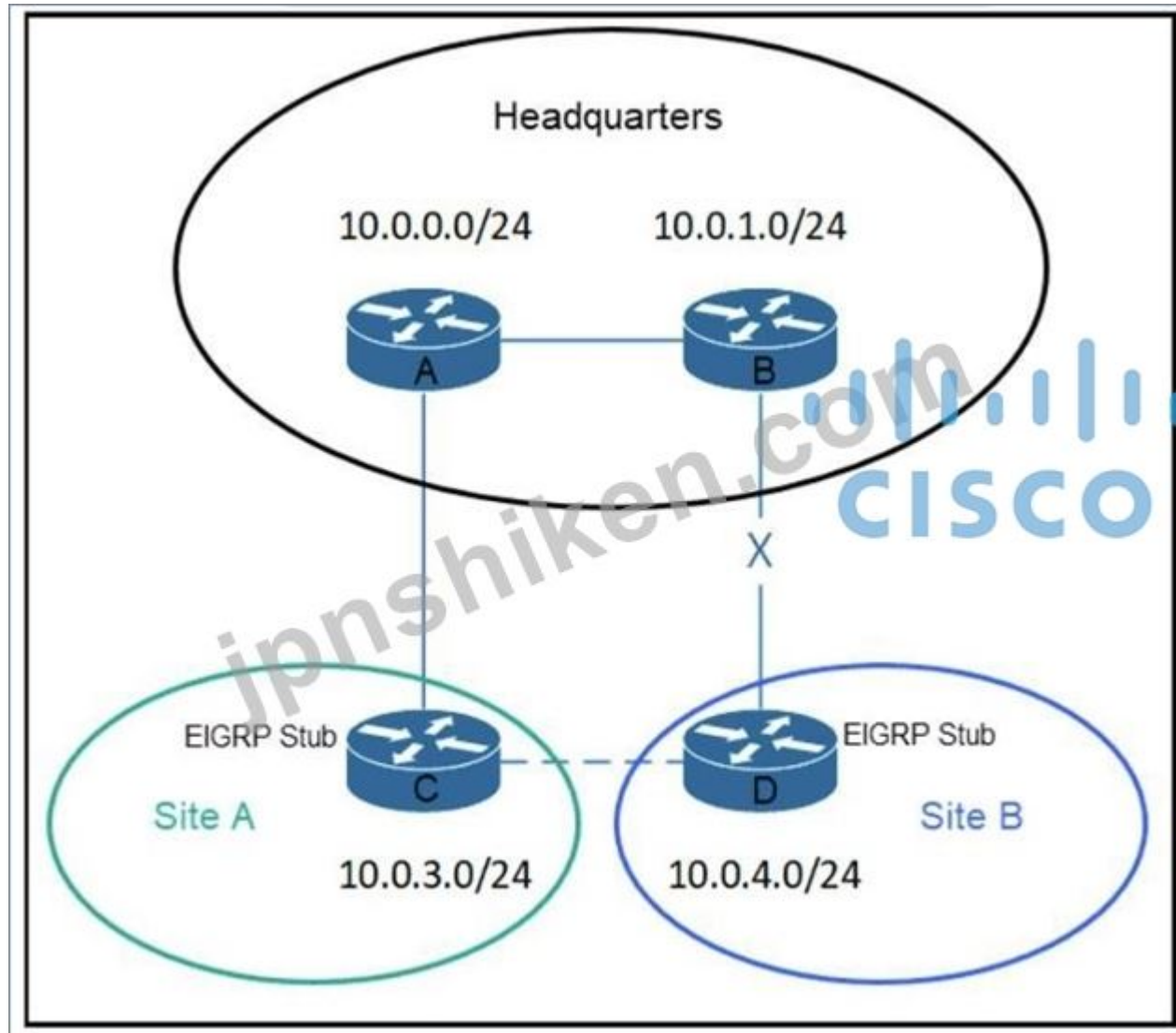
Explanation:

The WAN connectivity mapping distinguishes private, public, and overlay connectivity use cases. Enterprise WAN design normally selects the transport according to the required security, availability, latency, cost, and operational control. MPLS Layer 3 VPN is commonly used when the enterprise needs private routed WAN service with provider-managed segmentation and predictable QoS behavior. Internet VPN designs are used when cost, rapid deployment, or broadband availability is more important, with IPsec or SD-WAN providing encryption and overlay control. Layer 2 services such as VPWS or VPLS are selected when the customer requires Ethernet adjacency, VLAN transport, or transparent LAN-like connectivity across sites. The drag- and-drop answer preserves those distinctions by matching

each WAN type to the use case it is designed to support. A professional WAN architecture can combine these transports in a hybrid model, but the design must still identify which service carries private routing, which carries encrypted overlay traffic, and which provides Layer 2 extension. Reference topics: WAN transport selection, MPLS L3VPN, Internet VPN, Metro Ethernet, VPWS, VPLS, hybrid WAN design.

質問: 31

展示品を参照してください。



アーキテクトが会社用のルーティングソリューションを設計しています。新しい設計では、ルータ B と D 間のリンク障害時に 10.0.4.0/24 への接続が失われないようにするために、ルータ C と D に回線を追加します。アーキテクトはどのソリューションを選択する必要がありますか？

- A. スタブが接続されました
- B. スタブが再配布されました
- C. スタブ受信専用
- D. スタブリークマップ

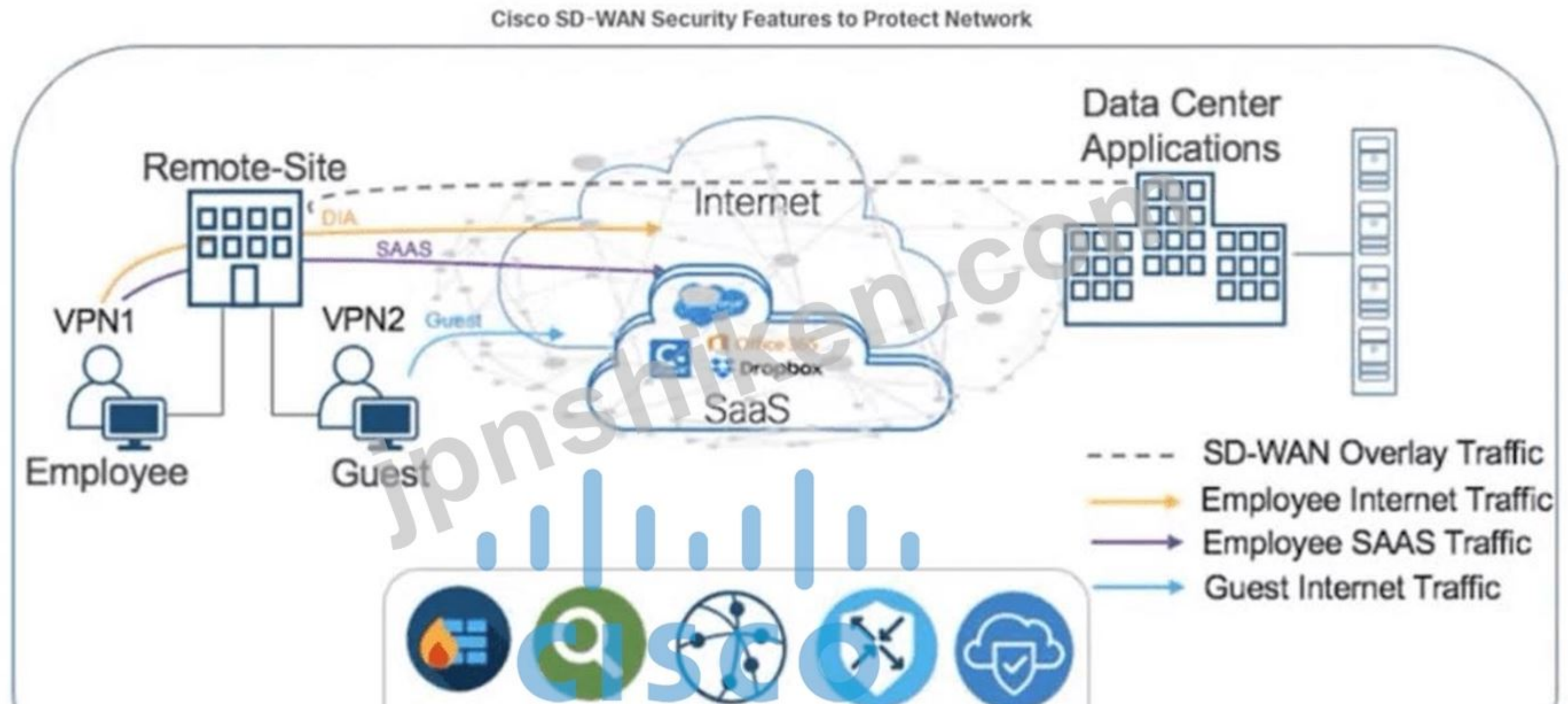
正解: [\(正解を表示します\)](#)

The EIGRP stub connected option is appropriate when the new circuit is intended to protect reachability to a connected subnet without turning the router into a broad transit point. EIGRP stub routing is used to limit the routes a router advertises to its neighbors and to constrain EIGRP query scope. A stub router can advertise connected routes, summary routes, redistributed routes, or combinations depending on the configured stub option. In this scenario, the new C-to-D circuit protects connectivity to 10.0.4.0/24 during a failure between B and D. Advertising the connected network through the stub behavior gives upstream routers a valid alternate route while still preventing unnecessary transit advertisements. Stub redistributed would be used for routes learned from another routing process, not for directly

connected networks. Stub receive-only would prevent the router from advertising any routes, which would not protect reachability. Stub leak-map is used for controlled exceptions and is unnecessary if the required route is a connected subnet. Reference topics: EIGRP stub routing, connected route advertisement, query boundary design, DUAL stability.

有効的な**300-420J**問題集はJPNTTest.com提供され、**300-420J**試験に合格することに役に立ちます！JPNTTest.comは今最新**300-420J**試験問題集を提供します。JPNTTest.com 300-420J試験問題集はもう更新されました。ここで**300-420J**問題集のテストエンジンを手に入れます。最新版のアクセス、<https://www.jpntest.com/shiken/300-420J-mondaishu> **457**問、**30%ディスカウント**、特別な割引コード：**JPNshiken**」

質問: 32



図を参照してください。Cisco Catalyst SD-WANのセキュリティ機能のうち、グローバルな脅威インテリジェンスと高度なサンドボックス機能を活用し、拡張ネットワーク全体でファイルアクティビティを継続的に分析する機能はどれですか？

- A. 侵入防止システム
- B. アプリケーション認識機能を備えたエンタープライズファイアウォール

C. Cisco Advanced Malware Protection

D. DNSレイヤーのセキュリティ

正解: (正解を表示します)

The described SD-WAN security feature is Cisco Advanced Malware Protection. Cisco Catalyst SD-WAN security documentation identifies AMP as the component that uses global threat intelligence, file reputation, advanced sandboxing, real-time malware blocking, and continuous analysis of file activity across the extended network. That wording maps directly to the question. AMP focuses on detecting malicious files, analyzing unknown files, blocking threats, and performing retrospective analysis if a file is later determined to be malicious. An intrusion prevention system inspects traffic for exploit signatures and known attack behaviors; it does not primarily provide file sandboxing and retrospective malware analysis. Enterprise Firewall with Application Awareness delivers stateful firewall control and application visibility. DNS-layer security, commonly integrated through Cisco Umbrella, blocks threats based on DNS requests before connections are made. The feature that specifically combines global threat intelligence, sandboxing, and continuous file analysis is AMP. Reference topics: Cisco Catalyst SD-WAN security, Advanced Malware Protection, UTD, file reputation, file analysis, sandboxing.

質問: 33

IS-ISプロトコルを使用してネットワークを設計する場合、どの2つの回線タイプがサポートされますか？ 2つ選択してください。）

A. 非ブロードキャストマルチアクセス

B. マルチアクセス

C. ポイントツーマルチポイント

D. 非放送

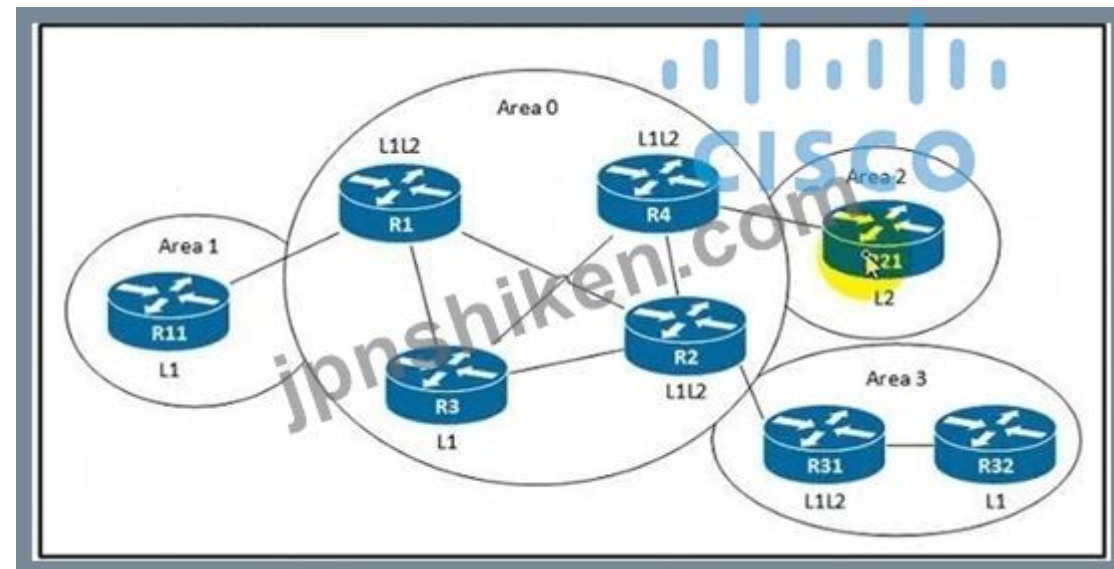
E. ポイントツーポイント

正解: B,E (コメントを发表する)

IS-IS supports two fundamental circuit types: broadcast multiaccess and point-to-point. A broadcast or multiaccess circuit is used on media such as Ethernet where multiple routers can share the same Layer 2 segment. In that case, IS-IS elects a designated intermediate system to simplify database synchronization on the shared network. A point-to-point circuit is used where exactly two IS-IS routers form an adjacency across a direct link, which is common in modern routed campus, data center, and WAN designs. IS-IS does not use the same nonbroadcast multiaccess and point-to-multipoint network-type model that is familiar from OSPF.

When IS-IS must run across technologies that do not naturally behave like Ethernet broadcast media, the design is typically adapted with point-to-point subinterfaces or appropriate Layer 2 treatment rather than selecting an IS-IS NBMA circuit type. Therefore, from the listed choices, the supported IS-IS circuit types are multiaccess and point-to-point.

質問: 34



図を参照してください。R1 と R2 間のリンクがダウンしたときに、顧客は予期しないネットワーク停止を経験しました。アーキテクトは、リンクが再度ダウンした場合にネットワークの継続性を確保するためのソリューションを設計する必要があります。

設計にはどのようなソリューションを含める必要がありますか？

- A. R31 を L1 ルーターにします。
- B. R3をL1L2ルーターにする
- C. エリア0をL2のみにする
- D. R11 を L2 ルーターにします。

正解: ([正解を表示します](#))

Making R3 an L1/L2 router restores the required IS-IS hierarchy and prevents the outage caused by the failed R1-to-R2 link. In IS-IS, Level 1 routers have detailed topology only inside their own area, while Level 2 routers provide the interarea backbone. Routers that connect a Level 1 area to the Level 2 backbone must operate as Level 1/Level 2 devices so they can participate in both databases and provide an exit path for area traffic. If a path between areas depends on a router that is only Level 1 or only Level 2 in the wrong location, traffic can be blackholed when a key link fails. Configuring R3 as L1/L2 gives the topology an additional valid transition point between the local area and the backbone, improving continuity after the R1-to-R2 failure. Making unrelated routers L1 or L2 does not necessarily create the required interarea attachment.

Reference topics: IS-IS Level 1, Level 2 backbone, L1/L2 routers, interarea reachability, hierarchical IGP design.

質問: 35

VRRP オブジェクト トラッキングに関する次の 2 つの記述のうち正しいものはどれですか。(2 つ選択してください)

- A. VRRPデバイスの優先度は、VRRPオブジェクトのアップまたはダウンステータスに応じて変更される可能性があります。
- B. VRRPインターフェースの優先度は管理者が手動で設定する必要があります
- C. VRRPグループは一度に1つのオブジェクトのみを追跡できます
- D. VRRPはインターフェースとルートのステータスを追跡できる
- E. VRRPはインターフェーストラッキングのみをサポートします

正解: ([正解を表示します](#))

VRRP object tracking allows the effective priority of a VRRP router to change when a tracked object changes state. Cisco documentation states that VRRP object tracking ensures the best VRRP router is master by altering VRRP priorities according to tracked objects such as interface or IP route states. This directly supports answer A. It also supports answer D, because VRRP can track interfaces and routes through the tracking process. The priority does not have to remain fixed; object tracking is specifically designed so a router can lower its priority when an uplink, route, or critical dependency fails. A VRRP group is not limited to one tracked object in modern Cisco implementations, so option C is too restrictive. VRRP does not support only interface tracking; route tracking is also a common use case, especially when the LAN interface is still up but upstream reachability has failed. In a resilient campus or branch gateway design, object tracking prevents a router from remaining master when it no longer has a valid upstream path, improving deterministic failover.

質問: 36



展示を参照してください。メール サーバー間の接続を提供するために、設計者はどの方法を使用する必要がありますか？

- A. ISATAP
- B. 6対4
- C. IPv4 コンパイル
- D. 6位

正解: ([正解を表示します](#))

The IPv4-compatible method is the closest match when the exhibit requires IPv6 connectivity between servers across an IPv4 portion of the network without using 6to4, 6rd, or ISATAP. IPv4-compatible IPv6 addressing represents an IPv4 address inside an IPv6 address format and was historically used for automatic IPv6-over-IPv4 tunneling between dual-stack nodes. Although this technique is now considered legacy and is not preferred for new designs, it matches the answer choice that directly embeds IPv4 reachability for host-to-host connectivity. ISATAP is normally used inside an enterprise IPv4 network to connect IPv6 hosts over an IPv4 intranet. 6to4 and 6rd are provider or site transition mechanisms that use specific IPv6 prefix behavior and are not simply a direct method between the depicted mail servers. For modern Cisco design, dual-stack or manually engineered tunnels would normally be preferred, but given these choices, IPv4-compatible addressing is the method aligned to the exhibit requirement. Reference topics: IPv6 transition mechanisms, IPv4-compatible addressing, tunneling, dual-stack migration, host connectivity.

質問: 37

エンジニアは、XML 表現で YANG を使用して、次の仕様に Cisco IOS XE スイッチを設定する必要があります。

直接接続されたホスト 10.10.10.1/27 からのインターフェイス GigabitEthernet2/1/0 接続で設定された IP アドレス 10.10.10.10/27 エンジニアが選択する必要がある YANG データ モデル セットはどれですか。

```
<interfaces xmlns="urn:ietf:params:xml:ns:yang:ietf-interfaces">
  <interface>
    <name>GigabitEthernet2/1/0</name>
    <type xmlns:ianaift="urn:ietf:params:xml:ns:yang:iana-if-type">ianaift:ethernetCsmacd</type>
    <enabled>false</enabled>
    <ipv4 xmlns="urn:ietf:params:xml:ns:yang:ietf-ip">
      <address>
        <ip>10.10.10.10</ip>
        <netmask>255.255.255.224</netmask>
      </address>
    </ipv4>
  </interface>
</interfaces>
```

A.

```
<interfaces YANG="urn:ietf:params:xml:ns:yang:ietf-interfaces">
  <interface>
    <name>GigabitEthernet2/1/0</name>
    <type YANG:ianaift="urn:ietf:params:xml:ns:yang:iana-if-type">ianaift:ethernetCsmacd</type>
    <enabled>true</enabled>
    <ipv4 YANG="urn:ietf:params:xml:ns:yang:ietf-ip">
      <address>
        <ip>10.10.10.10</ip>
        <netmask>255.255.255.224</netmask>
      </address>
    </ipv4>
  </interface>
```

B. </interfaces>

```

<interfaces json="urn:ietf:params:json:ns:yang:ietf-interfaces">
  <interface>
    <name>GigabitEthernet2/1/0</name>
    <type json:ianaift="urn:ietf:params:json:ns:yang:iana-if-type">ianaift:ethernetCsmacd</type>
    <enabled>true</enabled>
    <ipv4 json="urn:ietf:params:json:ns:yang:ietf-ip">
      <address>
        <ip>10.10.10.10</ip>
        <netmask>255.255.255.224</netmask>
      </address>
    </ipv4>
  </interface>
</interfaces>

```

C. </interfaces>

```

<interfaces xmlns="urn:ietf:params:xml:ns:yang:ietf-interfaces">
  <interface>
    <name>GigabitEthernet2/1/0</name>
    <type xmlns:ianaift="urn:ietf:params:xml:ns:yang:iana-if-type">ianaift:ethernetCsmacd</type>
    <enabled>true</enabled>
    <ipv4 xmlns="urn:ietf:params:xml:ns:yang:ietf-ip">
      <address>
        <ip>10.10.10.10</ip>
        <netmask>255.255.255.224</netmask>
      </address>
    </ipv4>
  </interface>
</interfaces>

```

D.

正解: ([正解を表示します](#))

The correct model set is the one that configures the IOS XE interface GigabitEthernet2/1/0 with IPv4 address

10.10.10.10 and prefix length 27 using valid XML according to the selected YANG schema. The directly connected host address 10.10.10.1/27 confirms the subnet relationship; both addresses belong to 10.10.10.0

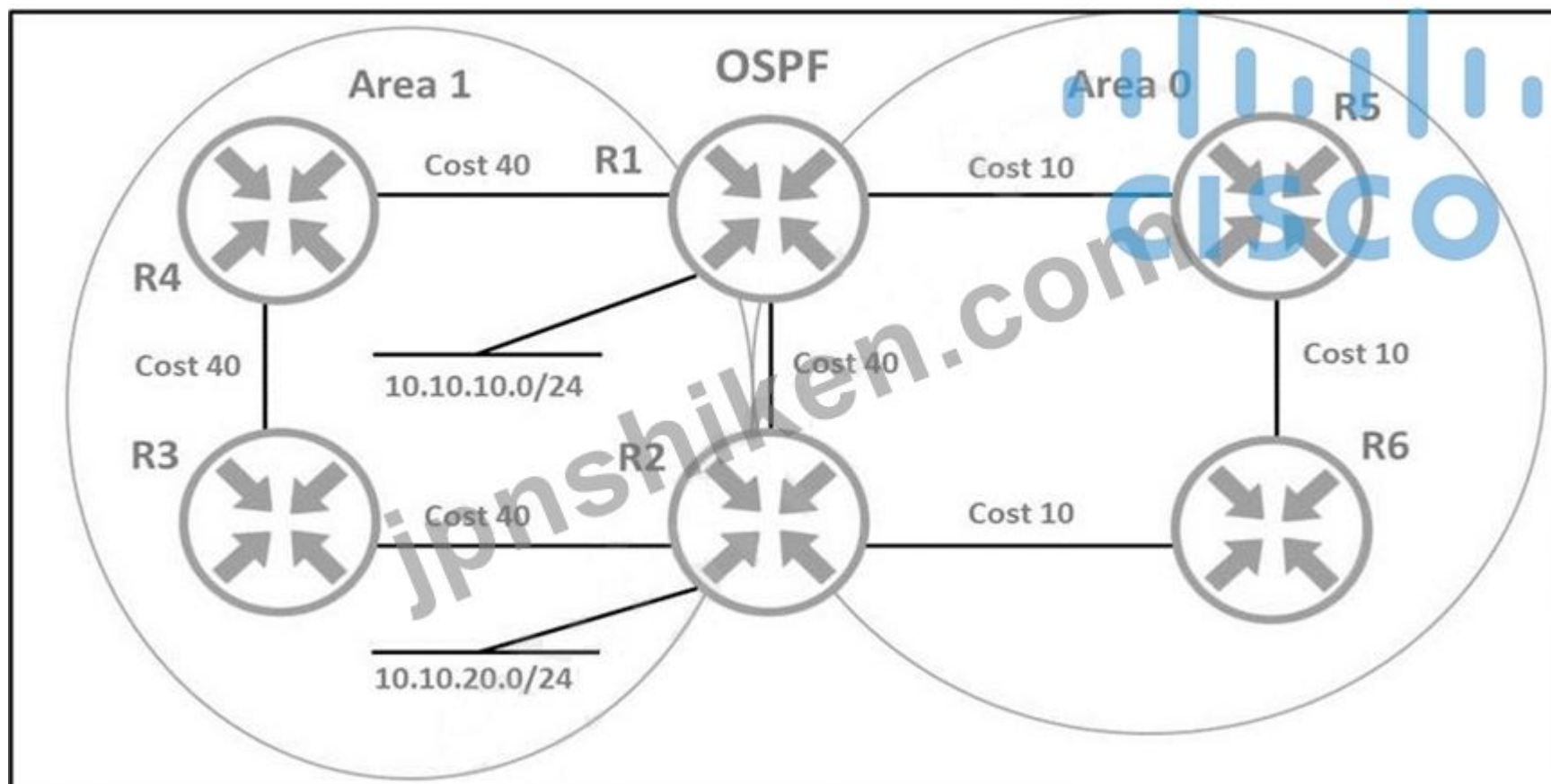
/27, so the interface configuration must use that prefix length rather than a host mask or an unrelated network.

YANG-based configuration is strict: the namespace, container hierarchy, interface type, interface name, and address fields must match the device model. If an option places GigabitEthernet under the wrong container, represents the mask incorrectly, or uses a schema from a different model family, the device will reject the payload or configure the wrong object. Option D correctly represents the IOS XE XML structure and the required addressing. In real operations, the engineer should follow this with a get-config readback against the running datastore to prove that the structured configuration committed as intended. Reference topics: IOS XE interface YANG, XML payload structure, NETCONF configuration, IPv4 subnet validation.

質問: 38

展示品を参照してください。

C0FD9 F48C9ACDC725EA850EC2476EE1E

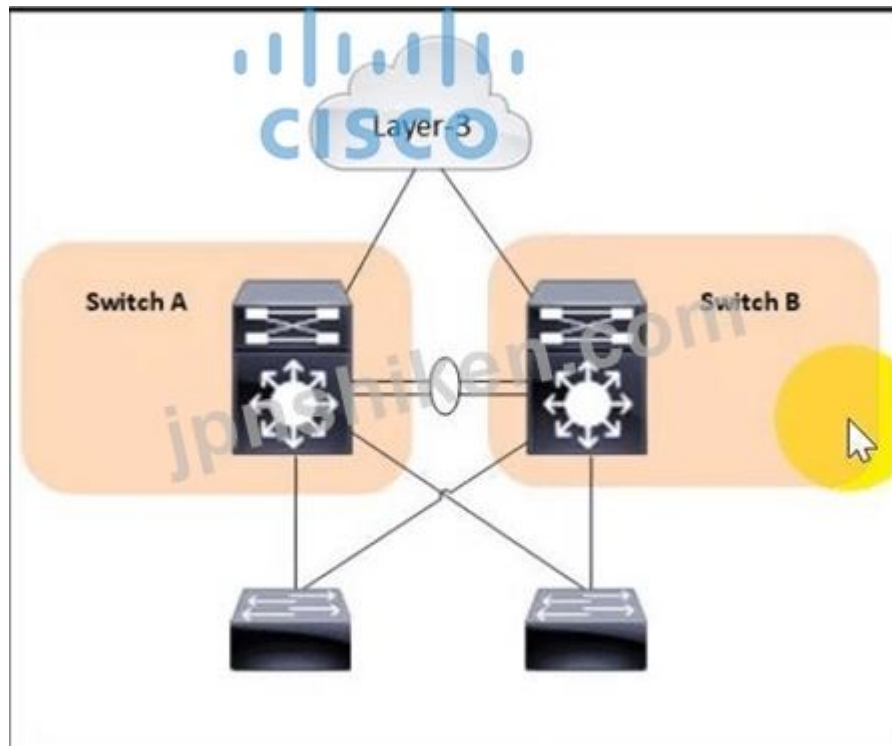


アーキテクトは、10.10.10.0/24 からネットワーク 10.10.20.0/24 へのトラフィックに R1 と R2 間の直接リンクを使用するソリューションを設計する必要があります。アーキテクトはどのソリューションを設計に含める必要がありますか。

- A. リンクの OSPF コストを 30 未満の値に設定します。
- B. OSPF エリア 0 の管理距離を下げます。
- C. リンクをエリア 2 に配置し、エリア 0 の R1 と R2 の間に新しいリンクをインストールします。
- D. マルチエリア隣接関係を提供するようにリンクを設定します。

正解: ([正解を表示します](#))

Lowering the OSPF cost on the direct R1-to-R2 link is the cleanest way to force traffic for 10.10.20.0/24 to prefer that link. OSPF calculates the shortest path using cumulative interface cost, and Cisco designs normally tune OSPF path selection by changing interface cost rather than administrative distance. Administrative distance determines preference between different routing protocols, not between OSPF paths inside the same OSPF process. Moving the link into another area or creating a multiarea adjacency adds design complexity and does not directly solve the requirement unless the resulting cost still becomes lower. The exhibit states that the architect must use the direct link for traffic from 10.10.10.0/24 toward 10.10.20.0/24. Therefore, the direct link must have a lower total OSPF metric than the alternate path. Configuring the link cost to a value below 30 makes OSPF choose it as the preferred intra-area route. This is deterministic, easy to validate with the OSPF database and routing table, and consistent with Cisco hierarchical IGP design practice. Reference topics: OSPF cost, SPF path selection, metric tuning, intra-area routing.



図を参照してください。顧客には、以下をサポートするように設計されたレイヤー 2 ネットワークが必要です。

- * 500個のアクティブ論理ポート
- * 30個のVLANのトランキング
- * 1秒未満の収束

どのスパニングツリー プロトコルを選択する必要がありますか？

- A. RPVST+
- B. MSTP
- C. CST
- D. PVST+

正解: ([正解を表示します](#))

MSTP is the correct spanning-tree choice for a Layer 2 design with many logical ports, multiple trunked VLANs, and a subsecond convergence requirement. Rapid PVST+ provides rapid convergence, but it creates a separate spanning-tree instance for every VLAN. With 30 VLANs on trunks and hundreds of active ports, the number of logical STP ports can grow quickly and increase CPU and memory load. MSTP maps many VLANs into a smaller number of spanning-tree instances while still using rapid convergence behavior based on IEEE 802.1w principles. This makes it more scalable for large Layer 2 domains than PVST+ or Rapid PVST+ when many VLANs are present. CST provides a single instance but lacks the flexibility and convergence characteristics expected in modern enterprise designs. The design needs both scalability and fast convergence, and MSTP is the standards-based protocol that best satisfies both requirements. Reference topics: MSTP, Rapid STP, logical STP ports, VLAN-to-instance mapping, Layer 2 scalability, convergence.

質問: 40

アーキテクトは、リモート オフィスを持つ企業がサポートする IPv6 移行ソリューションを設計する必要があります。

- * お客様は、サービス プロバイダーとの IPv4 ピアリングを使用しています。
- * IPv6 ユーザーは、IPv4 および IPv6 リソースにアクセスする必要があります。
- * 既存のコンテンツ プロバイダーは、今後 2 年間で IPv6 に移行します。
- * ユーザーは段階的に移行されます。

アーキテクトはどの移行ソリューションを選択する必要がありますか？

- A. NAT46

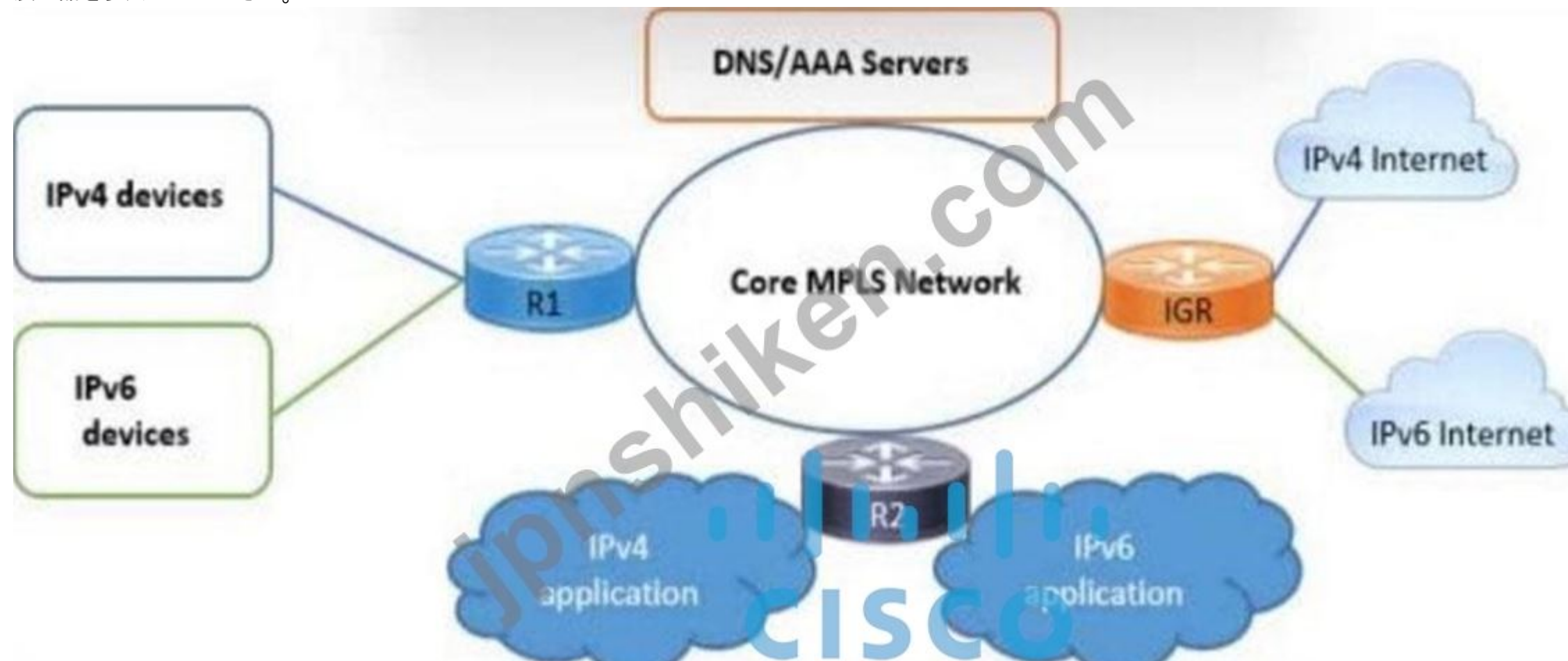
- B. トンネリング
- C. NAT64
- D. デュアルスタック

正解: ([正解を表示します](#))

Dual-stack is the correct IPv6 migration strategy for this corporation. The requirements state that users need access to both IPv4 and IPv6 resources, content providers will migrate over the next two years, and users will move in phases. Dual-stack supports exactly that model because hosts, routers, and services can run IPv4 and IPv6 simultaneously. Applications can select IPv6 or IPv4 based on DNS responses and protocol availability, while the organization gradually enables IPv6 without breaking existing IPv4 connectivity. NAT64 is useful for IPv6-only clients reaching IPv4-only services, but it introduces translation state and does not provide native dual protocol operation. NAT46 is a specialized translation direction and not the broad migration strategy here. Tunneling can connect IPv6 islands across IPv4 infrastructure, but the customer already has IPv4 peering and needs a phased endpoint and application migration, not only island interconnection. Dual-stack is straightforward operationally and gives the longest coexistence window while providers and applications transition. Reference topics: IPv6 migration, dual-stack deployment, DNS-based protocol selection, IPv4 and IPv6 coexistence, phased migration.

質問: 41

展示品を参照してください。



アーキテクトは、次の要件をサポートするために、企業顧客向けの IPv6 移行ソリューションを設計する必要があります。

- * クライアントは、IPv4 アドレスを指す同じ DNS 名を使用して、NAT64 および IPv6 DNS 解決サービスを提供する新しい IPv6 ネットワークに移行します。
- * サービス プロバイダーは、同じ IPv4 DNS サーバーを指す新しい IPv6 仮想アドレスを使用して、クライアント向けの IPv6 インターフェイスを作成します。
- * サービス プロバイダーは、グローバル IPv6 アドレスを使用し、IPv4 パケットを IPv6 トンネルにカプセル化するクライアントをサポートします。

アーキテクトはどの 2 つの移行ソリューションを選択する必要がありますか? (2 つ選択してください。)

- A. MPLS ネットワークから IGR へのデュアルスタック ライトを使用します。
- B. デバイスからコア MPLS ネットワークへの IPv6 トンネリングを使用します。

- C. デバイスからコア MPLS ネットワークまでデュアルスタック ライトを使用します。
- D. MPLS ネットワークから IGR への NAT44/64 を使用します。
- E. デバイスからコア MPLS ネットワークへの NAT44/64 を使用します。

正解: ([正解を表示します](#))

The design combines Dual-Stack Lite from the devices toward the core MPLS network and NAT44/64 from the MPLS network toward the Internet gateway router. The requirements describe IPv6-facing clients that still need access to IPv4 services and DNS behavior that allows the same service name to be reached during migration. Dual-Stack Lite carries IPv4 customer traffic across an IPv6 access or provider network by tunneling IPv4 packets over IPv6 toward a provider translation point. NAT44/64 services then provide translation or address-family interworking for IPv4-only resources. This combination is used when customer-facing access moves to IPv6 while legacy IPv4 services remain reachable during transition. A simple IPv6 tunnel from devices to the MPLS core does not provide the required IPv4 service translation. Using only NAT44/64 from end devices also shifts complexity to clients rather than placing translation at the network edge. The selected choices match the staged migration model: IPv6 transport with encapsulation in the access

/core side and translation at the provider or gateway boundary. Reference topics: IPv6 transition, DS-Lite, NAT64, DNS64, IPv4-over-IPv6 tunneling.

質問: 42

展示品を参照してください。



ネットワーク エンジニアは、以下に基づいてマルチキャスト ソリューションを設計する必要があります。

- * ユーザーとソース間の多対多の通信
- * 最大50のマルチキャストソースをサポート
- * Steamに登録する必要があるユーザー

エンジニアはどのマルチキャスト ソリューションを選択する必要がありますか？

- A. 任意のソースマルチキャスト
- B. 双方向PIM
- C. ソース固有のマルチキャスト
- D. マルチキャストVPN

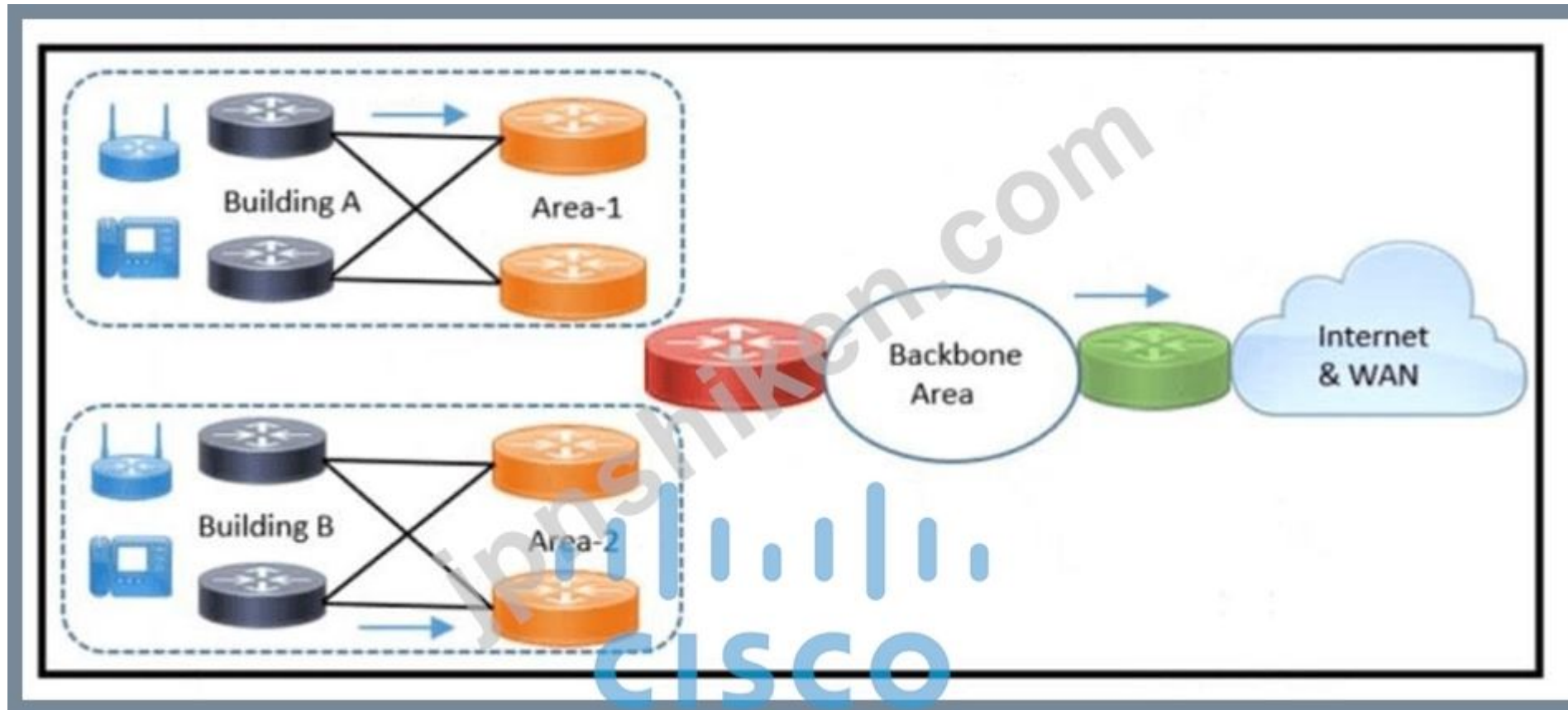
正解: B ([コメントを发表する](#))

Bidirectional PIM is the best fit for this multicast design because the application is many-to-many and has a manageable but meaningful number of multicast sources. Cisco describes BIDIR-PIM as a multicast mode designed for many-to-many applications where sources and receivers can be widely distributed. Unlike classic PIM sparse mode, BIDIR-PIM does not build source-specific shortest-path trees for each active source.

Instead, traffic uses a shared bidirectional tree rooted at the rendezvous point, which reduces multicast state in the network and improves scale when many sources may send to the same group. Source-Specific Multicast is optimized for one-to-many distribution where receivers know the source and join an (S,G) channel; it is not the natural many-to-many choice. Any Source Multicast can support many-to-

many traffic but usually creates more state because source trees may be built. Multicast VPN is a transport/service architecture, not the specific multicast mode requested. Reference topics: BIDIR-PIM, many-to-many multicast, shared trees, RP- based forwarding, multicast state scaling.

質問: 43



展示を参照してください。アーキテクトは、エンタープライズ顧客向けの OSPF ソリューションを設計する必要があります。設計は、次の要件を満たす必要があります。

リンク フラップの影響をエリア 1 とエリア 2 に制限します。

リンク障害が発生した場合でも、音声およびビデオ トラフィックへの影響は最小限に抑える必要があります。

アーキテクトが設計に含める必要がある 2 つの OSPF ソリューションはどれですか? (2 つ選択してください。)

- A. OR および BOR 選択の頻度を減らします。
- B. hello と how のタイマーを増やします。
- C. LSAおよびSPFスロットルタイマーを調整する
- D. 手動ルート集約を有効にし、すべての非バックボーン エリアをスタブ ネットワークとして設定します。
- E. バックボーンから非バックボーンエリアにデフォルトルートアドバタイズします。

正解: C,D ([コメントを发表する](#))

The OSPF design should include LSA/SPF timer tuning and summarization with appropriate stub-area design.

Cisco OSPF convergence is affected by failure detection, LSA generation, flooding, SPF scheduling, and routing-table installation. Tuning LSA and SPF throttling timers helps the network react quickly to genuine failures while avoiding excessive CPU churn during repeated flaps. This is especially important when voice and video traffic are sensitive to convergence delay. Manual route summarization at area boundaries reduces the number of detailed routes advertised between areas and limits the blast radius of link failures. Configuring nonbackbone areas as stub networks where appropriate also reduces external LSA propagation and lowers router resource use. Reducing DR and BDR election frequency is not a primary design tool for convergence, and increasing hello and hold timers would slow failure detection rather than improve it. Advertising default routes can be useful, but it does not by itself contain link flap impact. Reference topics: OSPF area design, LSA throttling, SPF throttling, route summarization, stub areas, convergence optimization.

質問: 44

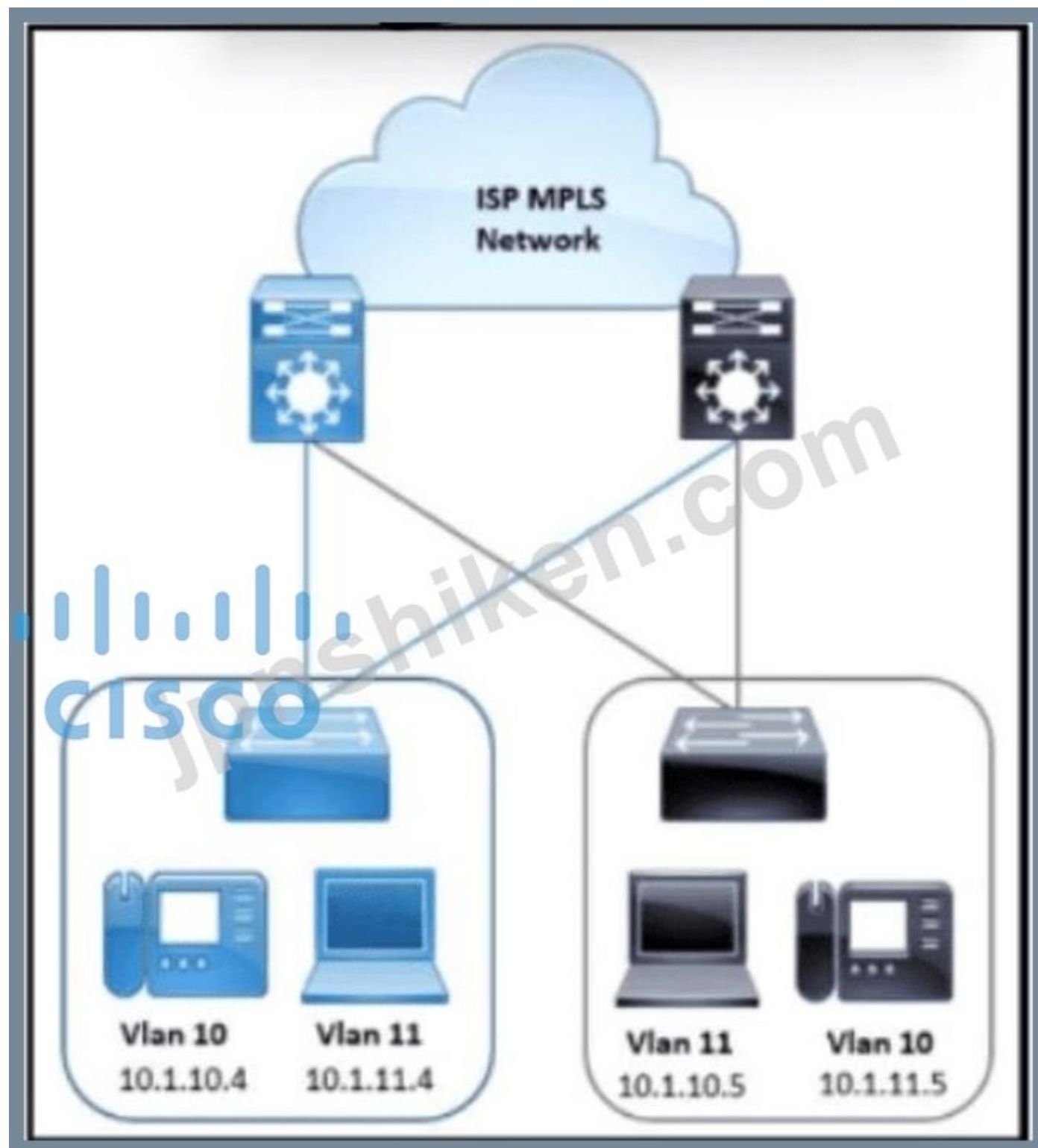
帯域外ネットワーク管理ソリューションを設計する利点は何ですか？

- A. 本番ネットワークが停止した場合でも、ネットワークデバイスを管理できます。
- B. 本番ネットワークと管理ネットワークの間に分離はありません。
- C. 本番ネットワークが停止した場合、バックアップネットワークパスとして使用できます。
- D. インバンド管理ソリューションよりも安価です

正解: A ([コメントを發表する](#))

The primary advantage of an out-of-band management network is that devices remain reachable for administration even when the production data network is impaired or unavailable. In an out-of-band design, management access uses a separate physical or logical path, often through dedicated management interfaces, console servers, terminal servers, or isolated management switches. This separation is valuable during routing failures, control-plane problems, misconfigurations, or security incidents affecting the production network. It also improves security because access control, authentication, logging, and monitoring can be concentrated on a management plane that is not directly exposed to ordinary user traffic. The statement that there is no separation from production traffic describes in-band management, not out-of-band management. An out-of- band network should not be treated as a general backup data path for production applications because that would defeat the isolation goal. It is also not necessarily less expensive; dedicated circuits, management switches, console servers, and operational processes can increase cost. Therefore, the correct benefit is continued manageability during a production network outage.

質問: 45



展示を参照してください。アーキテクトは、次の要件に基づいて、回復力のあるゲートウェイ ソリューションを設計する必要があります。

- * VLAN 10 および VLAN 11 は音声およびビデオ アプリケーションをサポートします。
- * リンクおよびノードの障害によるトラフィックへの影響は最小限に抑える必要があります。
- * 偽の hello パケットに対する保護を提供します。
- * IPv6をサポートします。

建築家はどのソリューションを選択する必要がありますか？

- A. IP SLA トラッキング付き GLBP
- B. 認証付き VRRP バージョン 2

C. MD5認証付きHSRPバージョン2

D. オブジェクト追跡機能付き VRRP バージョン 2

正解: **C** ([コメントを発表する](#))

HSRP version 2 with MD5 authentication is the correct resilient gateway choice because the design requires IPv6 support and protection against false hello packets. HSRPv2 supports both IPv4 and IPv6 gateway redundancy and improves group scaling compared with the original version. MD5 authentication protects the HSRP control exchange by ensuring that only authorized routers participate in the redundancy group, reducing the risk of a rogue or misconfigured device sending false hello messages and influencing active gateway selection. GLBP provides load balancing but does not satisfy the same IPv6 and authenticated hello requirement in this option set. VRRP version 2 is IPv4-focused and therefore does not meet the IPv6 requirement; VRRPv3 would be the relevant standards-based IPv6-capable version, but it is not offered.

Object tracking helps adjust gateway priority after link failure, but it does not protect against false hello packets by itself. Reference topics: HSRPv2, first-hop redundancy, IPv6 gateway redundancy, MD5 authentication, campus high availability.

質問: 46

ブランチ サイトにデュアル WAN エッジ ルーターを導入する場合、どのような設計上の考慮事項を考慮する必要がありますか？

A. BGP AS パスの先頭追加を使用して出力トラフィックに影響を与え、MED を使用してブランチからの入力トラフィックに影響を与えます。

B. HSRP の優先順位は、一方の WAN エッジを他方の WAN エッジよりも優先させる OMP ルーティング ポリシーと一致する必要があります。

C. DPI が適切に機能するには、トラフィックが WAN エッジから出てリモート サイトから戻るときに対称である必要があります。

D. 1 秒未満のリンク障害を検出するために、WAN エッジ ルータ間の BFD を設定します。

正解: ([正解を表示します](#))

When dual WAN Edge routers are deployed at a branch, first-hop redundancy and SD-WAN overlay path preference must be aligned. The correct design consideration is that HSRP priorities on the service side should match the OMP routing policy that determines which WAN Edge is preferred. If the LAN gateway points hosts to one WAN Edge while OMP policy prefers the other edge for return or remote-site traffic, traffic can become asymmetric or hairpinned through the wrong device. Proper alignment gives predictable active/standby or active/active behavior and avoids blackholing during failover. Option A describes traditional BGP path manipulation with AS-path prepending and MED; those are not the primary SD-WAN branch dual-edge design controls. Option C overstates symmetry as a universal requirement for DPI; Cisco SD-WAN can handle many traffic patterns, but gateway and OMP alignment remain the core branch design issue. Option D is wrong because SD-WAN already uses BFD across data-plane tunnels, not as a special direct adjacency between the two local WAN Edges for this purpose. Reference topics: Cisco SD-WAN branch redundancy, HSRP, OMP policy, WAN Edge preference, service-side high availability.

有効的な**300-420J**問題集はJPNTest.com提供され、**300-420J**試験に合格することに役に立ちます！JPNTest.comは今最新**300-420J**試験問題集を提供します。JPNTest.com 300-420J試験問題集はもう更新されました。ここで**300-420J**問題集のテストエンジンを手に入れます。最新版のアクセス、<https://www.jpntest.com/shiken/300-420J-mondaishu> **457**問、**30%ディスカウント**、特別な割引コード：**JPNshiken**」

質問: 47

インフラストラクチャ デバイスで同じ機能を設定する場合、IETF.OpenConfig と Cisco ネイティブ YANG モデルはどのように異なりますか？

A. OpenConfig モデルは IETF よりも包括的です。

B. Cisco ネイティブ モデルは OpenConfig ほど包括的ではありません。

C. Cisco ネイティブ モデルは IETF ほど包括的ではありません。

D. IETF モデルは OpenConfig よりも包括的です。

正解: ([正解を表示します](#))

OpenConfig, IETF, and Cisco native YANG models differ mainly in scope and implementation focus. IETF models are standards-based and are intended to cover common protocol or interface functions across vendors.

OpenConfig models are vendor-neutral operational and configuration models developed for broad network automation use cases and often provide a more detailed service or operational structure than the corresponding baseline IETF model. Cisco native models are device and operating-system specific, which means they usually expose the most complete Cisco feature coverage but are less portable across vendors. For the question wording, OpenConfig is more comprehensive than IETF when configuring the same general feature in a multivendor environment, because OpenConfig commonly defines richer operational and configuration containers beyond the minimum standards model. Option B is incorrect because Cisco native models are not generally less comprehensive than OpenConfig on Cisco devices. Option C is also incorrect for the same reason. Option D reverses the usual relationship. Reference topics: YANG model families, IETF models, OpenConfig models, Cisco native YANG models, model-driven programmability.

質問: 48

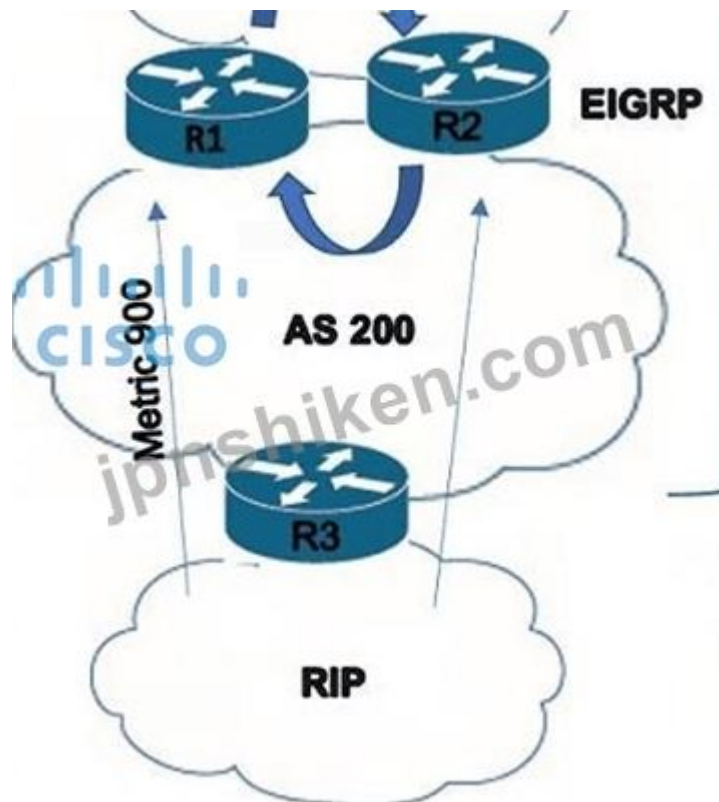
ある会社の支社では、本社への接続に冗長ルーターとリンクを使用しています。また、利用可能な帯域幅全体を使用するために、支社ではダイナミック ルーティング プロトコルを使用しています。設計者は、現在のネットワーク設計による RPF チェックの失敗を回避するために、マルチキャスト ストリーミング ソリューションを設計する必要があります。設計者はどの展開モデルを選択する必要がありますか？

- A. PIM-SM
- B. BIDIR-PIM
- C. PIM-BSR
- D. PIM-SSM

正解: ([正解を表示します](#))

BIDIR-PIM is the correct multicast deployment model when redundant paths and active links may create RPF challenges for many-to-many multicast traffic. Bidirectional PIM builds a shared tree rooted at the rendezvous point and avoids maintaining separate source-specific shortest-path trees for every sender. This is valuable when sources and receivers are spread across a redundant branch design and traffic can enter the multicast domain from multiple directions. Because BIDIR-PIM uses the Designated Forwarder mechanism on each link, it helps control which router forwards traffic toward the RP and reduces duplicate forwarding in topologies where normal RPF behavior could be problematic. PIM-SM is widely used, but it normally moves to source trees for active sources unless specifically controlled. PIM-SSM requires receivers to know the source and is not designed for arbitrary many-to-many communication. PIM-BSR is a mechanism for dynamic RP discovery, not a multicast forwarding model. Reference topics: BIDIR-PIM, multicast RPF, many-to-many multicast, Designated Forwarder, branch multicast design.

質問: 49



図を参照してください。アーキテクトは、R3 の背後にあるネットワークを EIGRP ネットワークに接続するソリューションを設計する必要があります。ルーティング ループを回避するには、どのメカニズムを含める必要がありますか。

- A. スプリットホライズン
- B. 要約
- C. ダウンビット
- D. ルートタグ

正解: ([正解を表示します](#))

Route tags are the correct loop-prevention mechanism when connecting another routing domain or redistributed network behind R3 into an EIGRP network. Redistribution can create route feedback: a route learned from one domain is redistributed into another, then later redistributed back into the original domain as if it were new information. That can produce routing loops or suboptimal path selection, especially in multipoint redistribution designs. Cisco route-map based redistribution commonly uses route tags to mark the origin of redistributed routes. Other redistribution points can then match those tags and prevent the same routes from being reintroduced into the source domain. Split horizon prevents advertisements back out the interface on which a route was learned, but it is not sufficient for multipoint redistribution loop prevention.

Summarization reduces route detail and query scope but does not identify route origin. The OSPF down bit is used in MPLS VPN/OSPF PE-CE scenarios, not general EIGRP redistribution. Therefore, the design should use route tags. Reference topics: EIGRP redistribution, route tagging, routing-loop prevention, route maps, redistribution policy.

質問: 50

エンジニアは Postman と YANG を使用して、ルーターを次のように構成します。

• OSPF process ID 100

• network 10.10.10.0/28 enabled for Area 0

どの get-config リプレイがモデル セットが正しく設計されたことを確認しますか？

```

<rpc-reply message-id="urn:uuid:1b3d05cd-8118-3e6a-6c05-012354678aaf" xmlns="urn:ietf:params:xml:ns:netconf:base:1.0" xmlns:nc="urn:ietf:params:json:ns:netconf:base:1.0">
  <data>
    <native json="http://cisco.com/ns/yang/ned/ios">
      <router>
        <ospf>
          <id>100</id>
          <network>
            <ip>10.10.10.0</ip>
            <mask>0.0.0.15</mask>
            <area>0</area>
          </network>
        </ospf>
      </router>
    </native>
  </data>
</rpc-reply>

```

```

<rpc-reply message-id="urn:uuid:1b3d05cd-8118-3e6a-6c05-012435678aaf" xmlns="urn:ietf:params:xml:ns:netconf:base:1.0" xmlns:nc="urn:ietf:params:xml:ns:netconf:base:1.0">
  <data>
    <native xmlns="http://cisco.com/ns/yang/ned/ios">
      <router>
        <ospf>
          <id>100</id>
          <network>
            <ip>10.10.10.0</ip>
            <mask>255.255.255.240</mask>
            <area>0</area>
          </network>
        </ospf>
      </router>
    </native>
  </data>
</rpc-reply>

```

```

<rpc-reply message-id="urn:uuid:1b3d05cd-8118-3e6a-6c05-021345678aaf" xmlns="urn:ietf:params:xml:ns:netconf:base:1.0" xmlns:nc="urn:ietf:params:xml:ns:netconf:base:1.0">
  <data>
    <native xmlns="http://cisco.com/ns/yang/ned/ios">
      <router>
        <ospf>
          <id>100</id>
          <network>
            <ip>10.10.10.0</ip>
            <mask>0.0.0.240</mask>
            <area>0</area>
          </network>
        </ospf>
      </router>
    </native>
  </data>
</rpc-reply>

```

```

<rpc-reply message-id="urn:uuid:1b3d05cd-8118-3e6a-6c05-394733226aaf" xmlns="urn:ietf:params:xml:ns:netconf:base:1.0" xmlns:nc="urn:ietf:params:xml:ns:netconf:base:1.0">
  <data>
    <native xmlns="http://cisco.com/ns/yang/ned/ios">
      <router>
        <ospf>
          <id>100</id>
          <network>
            <ip>10.10.10.0</ip>
            <mask>0.0.0.15</mask>
            <area>0</area>
          </network>
        </ospf>
      </router>
    </native>
  </data>
</rpc-reply>

```

- A. オプションA
- B. オプションB
- C. オプションC

D. オプションD

正解: ([正解を表示します](#))

The verifying get-config reply must reflect the same YANG model hierarchy and configured values that were pushed by the Postman request. In model-driven automation, a successful transport response is not enough; the engineer must confirm that the intended datastore contains the exact interface, routing, or protocol objects described by the model. Option D is the selected reply because it matches the expected YANG structure for the configuration operation shown in the exhibit. The key principle is that get-config returns configuration data from a specified datastore, such as running or candidate, using the schema-defined XML hierarchy. If the returned XML uses the wrong container, wrong namespace, incorrect key value, or an unexpected leaf, the model set was not designed or addressed correctly. Cisco NETCONF/YANG workflows rely on strict namespace and path accuracy, so even a small mismatch can result in configuration being placed in the wrong subtree or not applied. A clean validation checks the model namespace, parent containers, list keys, leaf names, and actual configured value. Reference topics: NETCONF get-config, YANG XML encoding, datastore validation, Postman automation workflow.

質問: 51

エンジニアは、金融アプリケーション用のマルチキャストネットワークを設計する必要があります。ほとんどのマルチキャストソースもマルチキャストトラフィックを受信します（多対多の展開モデル）。ルーティングテーブルをより適切にスケーリングするには、デザインでソースツリーを使用しないでください。これらの要件を満たすマルチキャストプロトコルはどれですか。

A. PIM-SSM

B. PIM-SM

C. MSDP

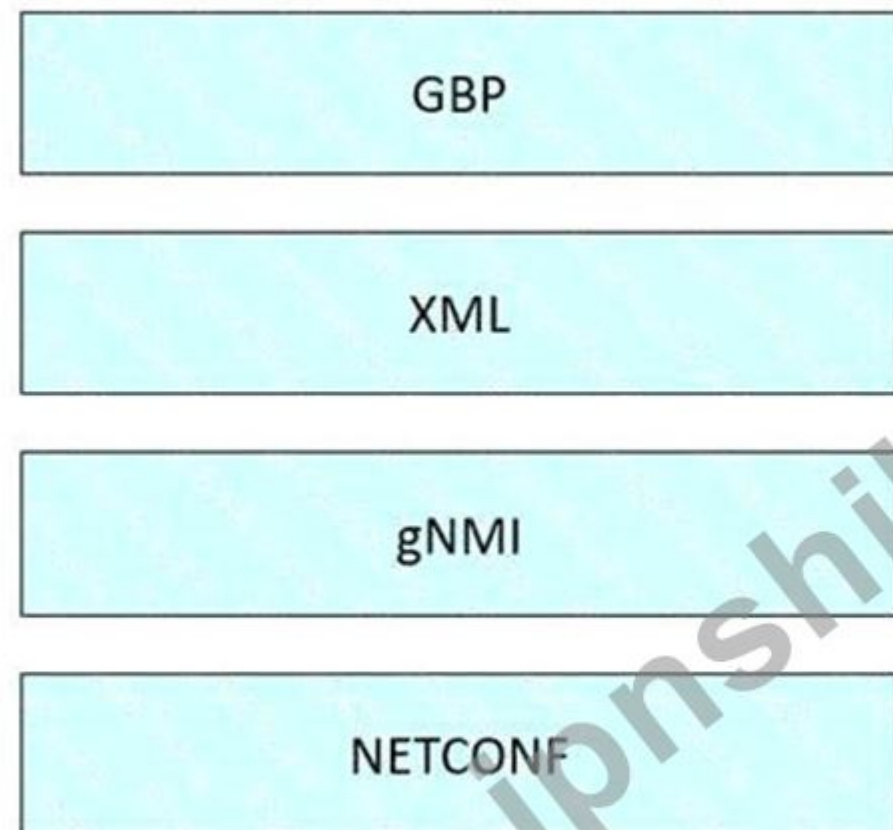
D. BIDIR-PIM

正解: D ([コメントを发表する](#))

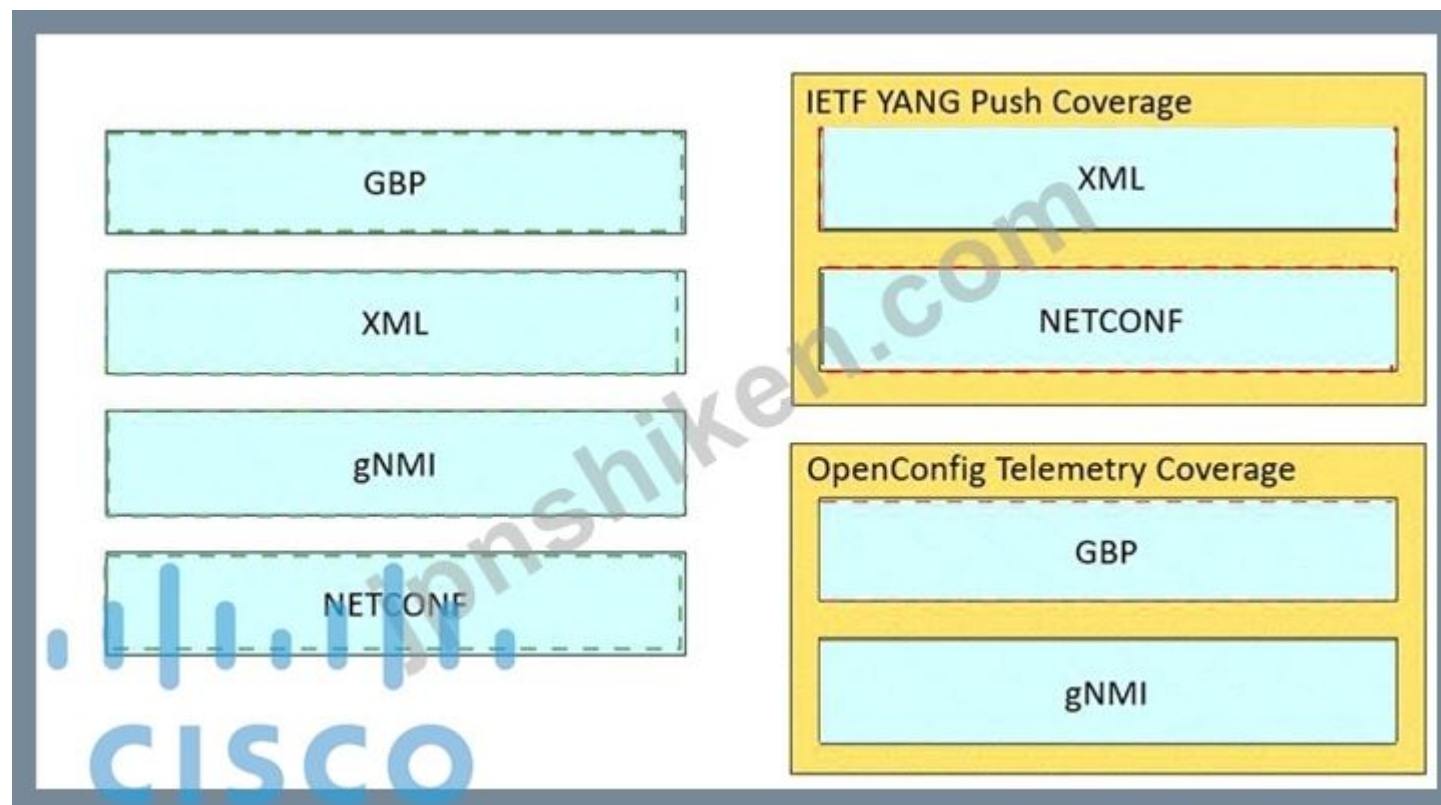
Bidirectional PIM is the multicast protocol mode designed for many-to-many applications where many receivers may also act as sources. In traditional PIM Sparse Mode, multicast forwarding can use shared trees through an RP and source-specific shortest-path trees. As the number of sources grows, maintaining source-specific state can increase routing table and control-plane scale requirements. BIDIR-PIM avoids source trees by using a shared bidirectional tree rooted at the rendezvous point address. Traffic from sources and traffic toward receivers uses the shared tree, which allows the multicast group to scale to a large number of sources without creating per-source state in the network. This aligns directly with the requirement that most multicast sources also receive multicast traffic and that source trees must not be used. PIM-SSM explicitly uses source-specific forwarding and requires receivers to know the source. PIM-SM may switch to source trees depending on behavior and policy. MSDP is used to exchange source information between PIM-SM domains, not to replace source trees inside one domain. Therefore, BIDIR-PIM is the correct design choice.

質問: 52

左側の要素を右側で使用されている YANG モデルにドラッグ アンド ドロップします。



正解:



Explanation:

The drag-and-drop answer follows the practical distinction between the YANG model families. IETF models are standards-based models intended to provide a common structure for widely implemented protocol and interface functions. OpenConfig models are vendor-neutral operational and configuration models developed for cross-vendor automation, with a strong focus on consistent telemetry and network state representation.

Cisco native models expose Cisco platform-specific capabilities and usually provide the deepest coverage for IOS XE, IOS XR, or NX-OS features that are not fully represented in a common model. The selected mapping therefore assigns model elements according to origin, portability, and feature depth. In a production automation design, engineers should start with open models when interoperability and long-term portability matter, but they should choose Cisco native models when a required Cisco feature is not available or is incomplete in IETF or OpenConfig. The choice is not cosmetic; it affects payload structure, namespace, validation behavior, and operational consistency across device families. Reference topics: YANG model families, IETF YANG, OpenConfig, Cisco native YANG, model-driven configuration.

質問: 53

1000を超えるルーターで構成されるEIGRPネットワークが設計されている場合、2つの有効なスケーリング技術は何ですか？ 2つ選択してください。）

- A. ルート要約を使用した構造化階層トポロジーを使用します
- B. 1秒未満のタイマーを使用
- C. ルートをフィルタリングするには、distribute-listコマンドを使用します
- D. リンクの遅延パラメーターを変更します
- E. 複数のEIGRP自律システムを実装する

正解: ([正解を表示します](#))

Large EIGRP networks require careful hierarchy and query-boundary design. A structured hierarchical topology with route summarization is one of the most important scaling techniques because summarization reduces routing table size, hides topology changes, and limits the EIGRP query scope when failures occur.

This improves convergence and reduces CPU and bandwidth consumption on routers that do not need detailed route visibility. Implementing multiple EIGRP autonomous systems can also be used as a scaling technique in very large designs because it creates administrative and query boundaries, although it introduces redistribution and policy complexity that must be controlled. Sub-second timers are not a primary scaling method; aggressive timers can increase control-plane load and do not reduce the size of the routing domain. Distribute lists can filter routes, but broad filtering alone is not the core scaling architecture for a network with more than

1000 routers. Modifying delay values influences metric calculation and path selection, not EIGRP domain scale. Therefore, the valid scaling techniques are structured hierarchy with summarization and, where justified, multiple EIGRP autonomous systems.

質問: 54



図を参照してください。ネットワーク管理者は、すべてのサブネットをアナウンスするのではなく、ルート集約を使用してサイトのサブネットを WAN にアナウンスする予定です。サイトのすべてのサブネットを網羅するために使用する必要がある最小の集約ルートは何ですか。

- A. 2001:DB8:ABCD:0003::/60
- B. 2001:DB8::732
- C. 2001:DB8:ABCD::760
- D. 2001 DB8 ABCD /64

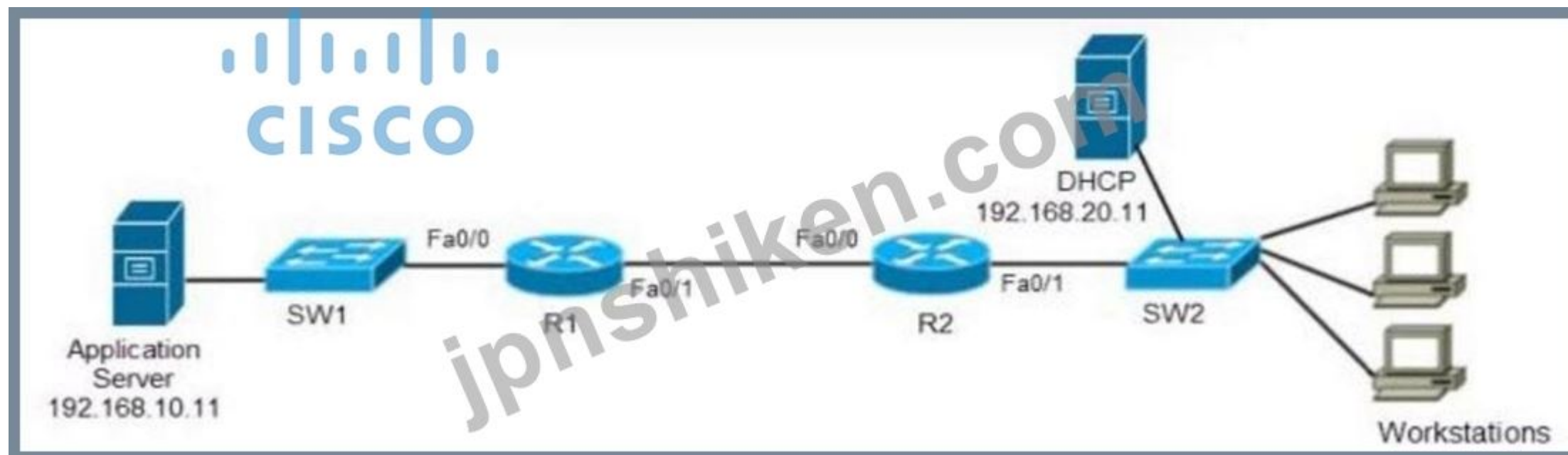
正解: C ([コメントを发表する](#))

The selected IPv6 summary must be the smallest prefix that covers all of the site subnets in the exhibit while not including unnecessary additional networks. IPv6 route summarization works by identifying the common high-order bits shared by the component prefixes and advertising the shortest aggregate that contains them all.

Cisco routing design uses summarization at WAN or aggregation boundaries to reduce routing-table size, contain convergence events, and simplify policy. Option C is retained because it represents the aggregate identified by the exhibit answer set. Option A is a different /60 boundary and would be correct only if all component subnets fell under that exact nibble range. Option B is not a usable summary route in normal IPv6 prefix notation as written. Option D is too broad or malformed for a precise site summary. The professional design rule is to summarize on valid nibble or bit boundaries and verify that every component prefix is included without accidentally absorbing unrelated site networks. Reference topics: IPv6 summarization, prefix aggregation, route boundary calculation, WAN route scale, convergence containment.

質問: 55

展示品を参照してください。



アーキテクトは、Wake-on-LAN アプリケーションをサポートする顧客向けのネットワークを設計しています。アーキテクトはどのソリューションを選択する必要がありますか？

- A. R1 上の IP ダイレクトブロードキャスト
- B. SW1 上のスパニングツリー アップリンクファースト
- C. SW2 上のスパニングツリー アップリンクファースト
- D. R2 上の IP ダイレクトブロードキャスト

正解: ([正解を表示します](#))

IP directed broadcast must be enabled on the last Layer 3 device connected to the destination client subnet, which is R2 in the exhibit. Wake-on-LAN uses a magic packet that must reach sleeping hosts on the target LAN. Because sleeping clients often do not have an active IP stack or ARP entry, the WoL packet is commonly sent as a directed broadcast toward the target subnet. Routers forward the packet as a unicast until it reaches the destination subnet, where the last-hop router converts it into a Layer 2 broadcast if directed broadcasts are permitted. Cisco disables directed broadcasts by default on many platforms because of historical amplification-attack risks, so the design should enable it only where needed and protect it with ACLs. Spanning-tree UplinkFast is unrelated to delivering WoL packets; it improves STP convergence on access switches. Enabling directed broadcasts on the wrong router would not place the magic packet onto the client VLAN. Therefore, R2, the last-hop router toward the target WoL clients, is the correct location.

Reference topics: Wake-on-LAN, IP directed broadcast, last-hop router behavior, subnet broadcast forwarding, ACL protection.

質問: 56

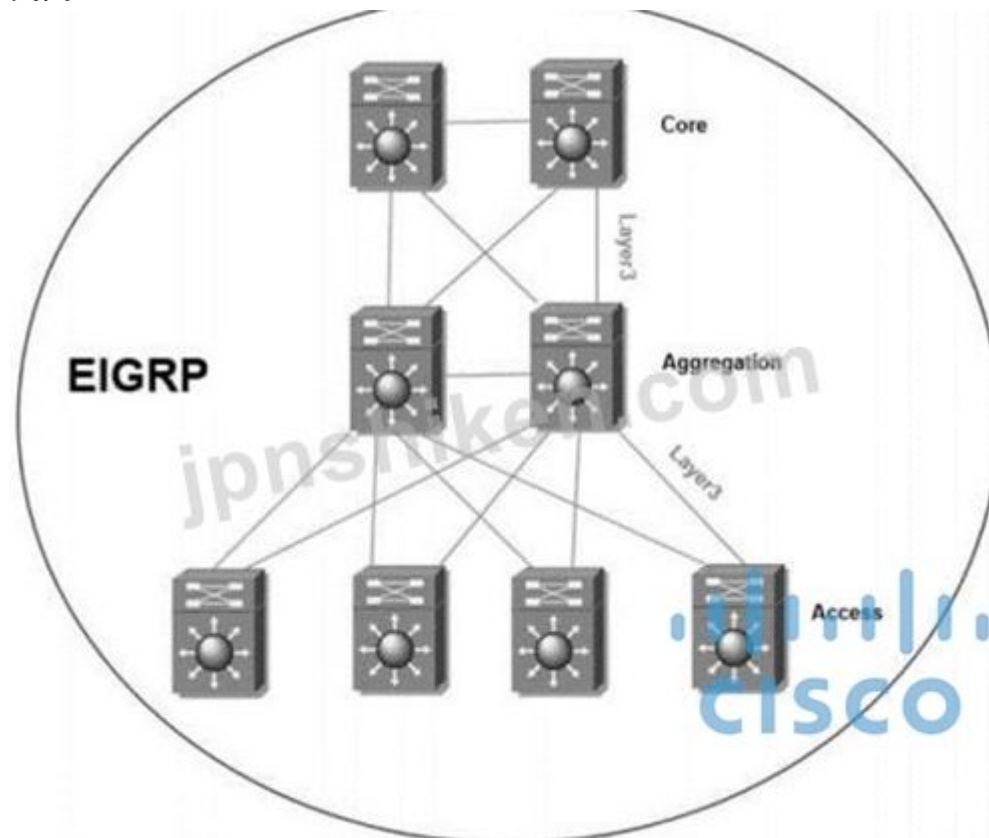
一般的な Cisco SD-WAN 展開でフルメッシュ IPsec トンネルを確立する前に、WAN エッジ ルータはデータ プレーン暗号化のキー情報を交換するためにどのメカニズムを使用しますか。

- A. キー交換サーバーとして vSmart コントローラーを使用します。
- B. キー交換サーバーとして vManage を使用します。
- C. 相互にキーを交換するときに IKEv2 を使用します。
- D. キー交換サーバーとして vBond を使用します。

正解: ([正解を表示します](#))

Cisco SD-WAN uses the vSmart controller as the control-plane component that distributes the information needed for secure data-plane tunnel establishment. Before WAN Edge devices build full-mesh IPsec tunnels, they must learn reachability and security information through the SD-WAN control plane. Cisco documentation describes the vSmart controller as the centralized control-plane element that exchanges OMP routes, policy, TLOC information, and security parameters with WAN Edge routers. In the traditional SD-WAN key-exchange model, vSmart sends IPsec encryption keys to edge devices so the data plane can be secured without each edge independently negotiating IKEv2 with every other edge. vManage is the management platform and is not the key exchange point. vBond performs orchestration, authentication assistance, and NAT traversal during onboarding, but it is not the data-plane key server. IKEv2 may be used in specific newer or configured security models, but the typical Cisco SD-WAN overlay relies on vSmart for scalable key distribution. Reference topics: Cisco SD-WAN control plane, vSmart, OMP, IPsec key distribution, WAN Edge security.

質問: 57



展示を参照してください。完全な EIGRP ルーティング テーブルがネットワーク全体にアドバタイズされます。現在、ネットワーク内の 1 つのリンクに障害が発生すると、ユーザーはデータ損失を経験します。アーキテクトは、リンクに障害が発生したときの影響を軽減するためにネットワークを最適化します。アーキテクトはどのソリューションを設計に含める必要がありますか？

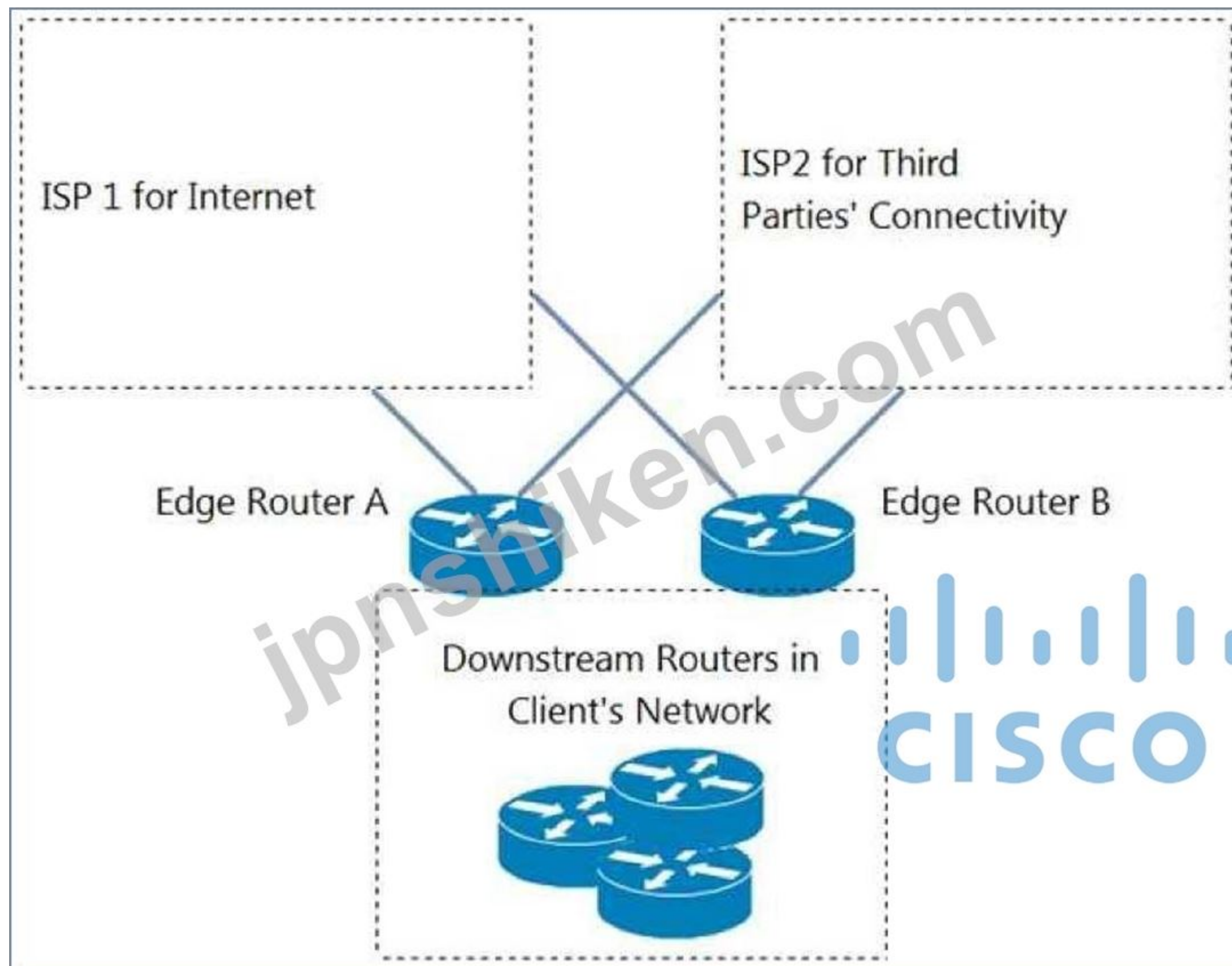
- A. EIGRP ネイバー間のインターリンクで BFD を実行します。
- B. 各アクセス レイヤー スイッチから集約レイヤーへのアクセス レイヤー ネットワークを要約します。
- C. デフォルトの EIGRP hello 間隔とホールド時間を短縮します。
- D. 集約層からコア層に向かってアクセス層ネットワークを要約します。

正解: A ([コメントを发表する](#))

Running BFD between EIGRP neighbors is the appropriate design when the problem is traffic loss during link failure and the requirement is to reduce the impact of the failure. EIGRP can detect neighbor loss through its hello and hold timers, but those timers may not be fast enough for applications sensitive to packet loss.

Bidirectional Forwarding Detection provides rapid failure detection independent of the routing protocol's normal timer expiration. When BFD detects a failed forwarding path, it notifies EIGRP so the routing process can converge without waiting for the EIGRP hold timer. Summarization is still a valuable scaling technique, but it primarily reduces route table size and query scope; it does not directly accelerate failure detection on an individual inter-router link. Reducing EIGRP hello and hold timers can improve detection, but it increases control-plane load and is less efficient than BFD. The question emphasizes data loss when any one link fails, so fast failure detection is the best corrective design. Reference topics: EIGRP convergence, BFD, failure detection, hello and hold timers, routed-link resiliency.

質問: 58



図を参照してください。エンジニアは、完全なインターネット接続のために ISP1 とピアリングし、複数のサードパーティのルートを直接交換するために ISP2 とピアリングするクライアント用の BGP ソリューションを設計しています。エッジ ルーターに実装すると、どのアクションによってクライアント ネットワークが ISP1 を介してインターネットに到達できるようになりますか。

- A. 各 ISP の異なる VRF 内で eBGP セッションを実行します。
- B. クライアント ネットワーク内のダウンストリーム ルーターのデフォルト ルートをアドバタイズします。
- C. ISP2 に AS パス プリペンド機能を適用します。

D. クライアントが自身の AS から発信されたルートのみをアドバタイズするようにルート フィルタリングを適用します。

正解: ([正解を表示します](#))

The client edge can reach the Internet through ISP1 only if downstream routers inside the client network have a usable route toward the Internet. Since ISP1 provides full Internet connectivity and ISP2 is used for direct exchange with selected third parties, the edge design should advertise a default route into the client network.

This gives internal routers a simple, scalable route for all Internet destinations while preserving specific routing where necessary. Running separate VRFs for each ISP may be a valid segmentation technique, but it does not by itself tell downstream routers how to reach the Internet. AS-path prepending toward ISP2 influences inbound path preference from external networks; it does not solve internal default routing. Filtering outbound advertisements so the client advertises only its own originated routes is important to avoid becoming a transit AS, but it is not the mechanism that enables internal users to reach the Internet through ISP1. The clean design is to receive or originate a default route at the edge and advertise that default to the client routing domain. Reference topics: BGP Internet edge design, default route advertisement, transit prevention, ISP peering policy.

質問: **59**

同じ2つのBGP機能のうち、同じAS番号を共有するeBGPネイバー間のルート交換が成功するのはどれですか。 2つ選択してください。)

- A.** advertise-best-external
- B.** bestpath as-path ignore
- C.** client-to-client reflection
- D.** as-override
- E.** allow-as-in

正解: ([正解を表示します](#))

When eBGP neighbors share the same autonomous system number, normal BGP loop-prevention behavior can block route acceptance because the receiving router sees its own AS in the AS_PATH. Two BGP features address this design condition. The allow-as-in feature permits a BGP router to accept routes even when its own AS appears in the AS_PATH a configured number of times. This is often used in dual-homed or migration designs where the same customer AS appears at multiple sites. The as-override feature is typically used by a provider edge router to replace the customer AS in the AS_PATH with the provider AS before advertising the route to another customer site that uses the same AS. Together, these features allow successful route exchange in same-AS eBGP scenarios while preserving loop-prevention intent through controlled policy. Advertise-best-external influences advertisement of best external paths and does not solve same-AS rejection. Bestpath as-path ignore affects path selection, not route acceptance. Client-to-client reflection is an iBGP route-reflector behavior, not an eBGP same-AS solution.

質問: **60**

組織は、2つの異なる自律システムにマルチキャストを展開することを計画しています。彼らのソリューションでは、RPが次のことを行えるようにする必要があります。

*ドメイン外のアクティブなソースを発見する

*他のRPとの接続に基礎となるルーティング情報を使用する

*グループに参加しているソースを発表する

これらの要件をサポートするソリューションはどれですか？

- A.** MSDP
- B.** SSM
- C.** PIM-SM
- D.** PIM-DM

正解: ([正解を表示します](#))

MSDP is the correct solution for multicast across different autonomous systems when RPs must discover active sources outside their local domain. In PIM Sparse Mode, each multicast domain can have its own rendezvous point. Without a mechanism to exchange source information, receivers in one domain may not know about active sources registered with an RP in another domain. Multicast Source Discovery Protocol solves that problem by allowing RPs to establish peer relationships and exchange Source-Active information.

Cisco multicast design uses MSDP with PIM-SM when separate domains need to share multicast source reachability while relying on the unicast routing table for connectivity between RPs. SSM does not use an RP and depends on receivers joining a specific source and group. PIM-SM by itself works inside a domain but does not provide interdomain source discovery between RPs. PIM-DM is flood-and-

prune and is not suitable for this sparse, interdomain design. Therefore, the organization should deploy MSDP between the RP devices or RP domains. Reference topics: MSDP, PIM-SM, interdomain multicast, Source-Active messages, RP-to- RP source discovery.

質問: 61

図を参照してください。メディアコンテンツを制作する中規模企業は、国内に4つのオフィスを構え、ローカルISPが提供するMPLSレイヤ3サービスで接続しています。ネットワークはスタティックルーティングを使用しています。将来の成長を見越して、エンジニアリングチームはRFC 5340の要件に従って設計の改善を検討し、推奨する必要があります。ソリューションはルーティングテーブルを最適化し、ルータ間で交換されるルーティングアップデートの数を削減する必要があります。更新されたルーティング設計は信頼性が高く、ルーティングループを回避する必要があります。どの実装が要件を満たしていますか？

- A. 複数のAS番号を使用したEIGRP
- B. スタブエリアを使用したOSPF
- C. 各ロケーションごとに固有のアドレスファミリーを持つBGP
- D. スタブエリアルータを備えたOMP

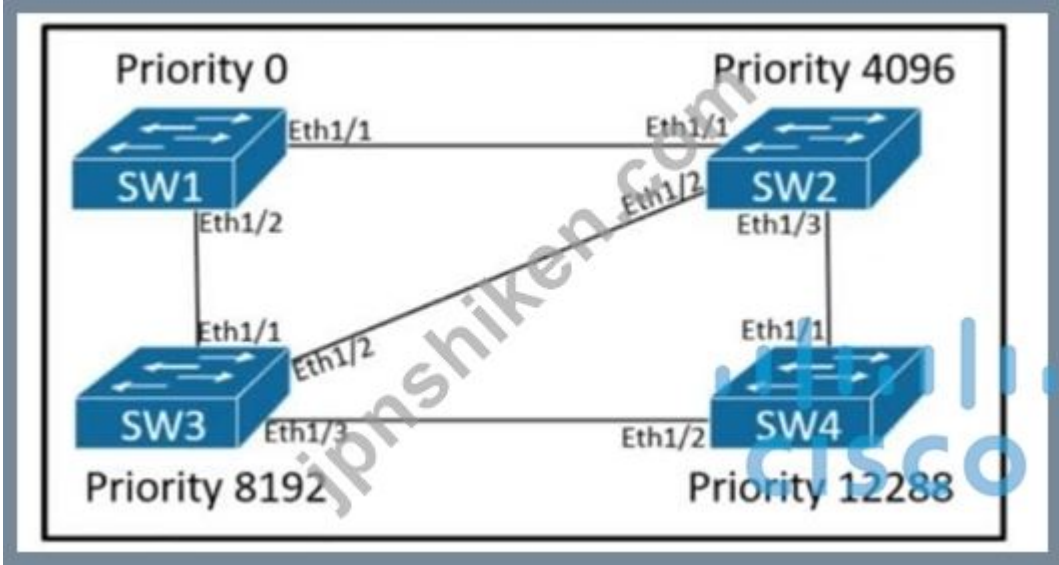
正解: (正解を表示します)

OSPF with stub areas is the correct implementation. The reference to RFC 5340 points to OSPFv3, which Cisco documents as the OSPF version used for IPv6 and also capable of carrying IPv4 unicast address families in modern implementations. OSPF is a link-state routing protocol that provides loop-free convergence through SPF calculation, making it a stronger fit than static routing or RIP for a growing WAN.

Stub areas reduce routing information inside branch or regional areas by blocking external LSAs and allowing the ABR to inject a default route. That reduces route-table size and limits update propagation, which matches the requirement to optimize routing tables and reduce routing updates between routers. EIGRP with multiple autonomous systems is not aligned with RFC 5340. BGP address families are useful in provider and MPLS VPN designs but are not the best internal route-optimization answer here. OMP is specific to Cisco SD-WAN, not a generic RFC 5340 routing design. Reference topics: OSPFv3, RFC 5340, stub areas, LSA reduction, SPF convergence.

有効的な**300-420J**問題集はJPNTTest.com提供され、**300-420J**試験に合格することに役に立ちます！JPNTTest.comは今最新**300-420J**試験問題集を提供します。JPNTTest.com 300-420J試験問題集はもう更新されました。ここで**300-420J**問題集のテストエンジンを手に入れます。最新版のアクセス、<https://www.jpntest.com/shiken/300-420J-mondaishu> **457**問、**30%ディスカウント**、特別な割引コード：**JPNshiken**」

質問: 62



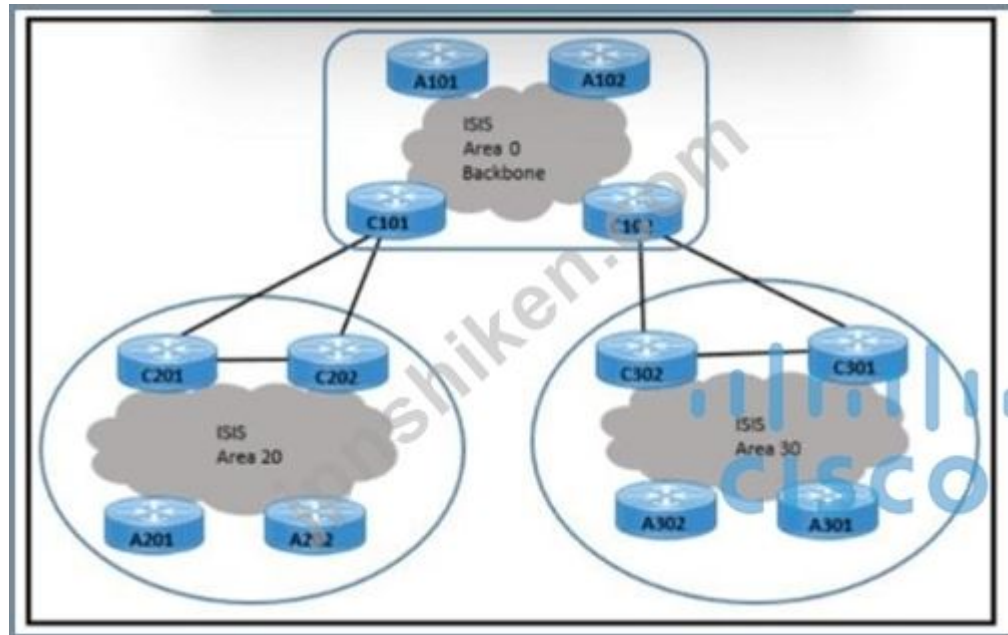
展示を参照してください。エンジニアは、ループのないレイヤー 2 ネットワーク設計を必要とする企業にこのソリューションを提案しました。ネットワークは 802.1W を実行し、すべてのリンクは 1 Gbps になります。すべてのインターフェイスがポイントツーポイント隣接関係として稼働している場合、設計に基づいて予想されるポート終了状態は何ですか？

- A. SW2とSW3のEth1/2はDesg FWD状態になります
- B. SW2とSW3のEth1/3はAttn BLK状態になります
- C. SW3 および SW4 の Eth1/2 は Attn BLKbtate になります。
- D. SW1 および SW2 上の Eth1/1 は Root FWD 状態になります。

正解: [C \(コメントを发表する\)](#)

The expected port state follows Rapid Spanning Tree Protocol role selection on equal-speed point-to-point links. With 802.1w, each non-root switch selects one root port based on the lowest path cost to the root bridge. Where there are redundant Layer 2 paths of equal speed, one side of a redundant segment must remain blocked as an alternate port to maintain a loop-free topology. Because all links are 1 Gbps, path cost is equal unless bridge priority, bridge ID, or port ID breaks the tie. The selected option identifies the port pair that becomes alternate blocking based on the root placement and tie-breakers shown in the exhibit. RSTP improves convergence compared with classic 802.1D by using proposal and agreement handshakes on point-to-point links, but it still blocks redundant paths when a pure Layer 2 loop exists. A design requiring all links to forward would need a routed access design, MEC, StackWise Virtual, VSS, or another loop-free multipath technology. Reference topics: RSTP, root port, designated port, alternate port, point-to-point links, Layer 2 loop prevention.

質問: 63



展示を参照してください。アーキテクトは、次のような要件を持つ顧客向けに階層型 IS-IS ソリューションを設計しています。

- * 今後 24 か月以内に、すべての地域でルーターの数が 2 倍になります。
- * エリア 20 および 30 内のリンク フラップは、バックボーン エリアに影響を与えてはなりません。
- * A201 および A302 ルータから発信されるトラフィックは、バックボーンのアプリケーション サーバーに接続する必要があります。

建築家はどのデザインを選択する必要がありますか？

- A. C201 レベル 1/2、A301 レベル 1/2、A102 レベル 1/2
- B. C101 レベル 1/2。A201 レベル 1、および A101 レベル 2
- C. C102 レベル 2、A202 レベル 2、および A102 レベル 1
- D. C302 レベル 2。A302 レベル 1/2。および A101 レベル 2

正解: [\(正解を表示します\)](#)

The selected IS-IS design must place the correct routers as Level 1/Level 2 boundary routers so that area failures stay local while traffic can still reach backbone application services. In a hierarchical IS-IS design, Level 1 routers exchange detailed topology only inside their own area. Level 2 routers form the backbone used for interarea connectivity. L1/L2 routers connect the local area to the backbone and

advertise reachability between the two levels. Areas 20 and 30 are expected to double in size, and link flaps inside those areas must not impact the backbone. Therefore, the design should keep internal area routers at Level 1 and use only the boundary devices that connect those areas to the backbone as L1/L2. Option A matches that principle by selecting the appropriate boundary routers for interlevel connectivity while avoiding unnecessary L1/L2 adjacencies throughout the topology. The other options either place the wrong routers in Level 2 only or fail to provide a clean transition between local areas and the backbone. Reference topics: IS-IS hierarchy, L1/L2 boundary design, area containment, backbone stability, route leaking.

質問: 64

オーバーレイ管理プロトコルがSD-WANオーバーレイでアドバタイズするルートはどれですか。

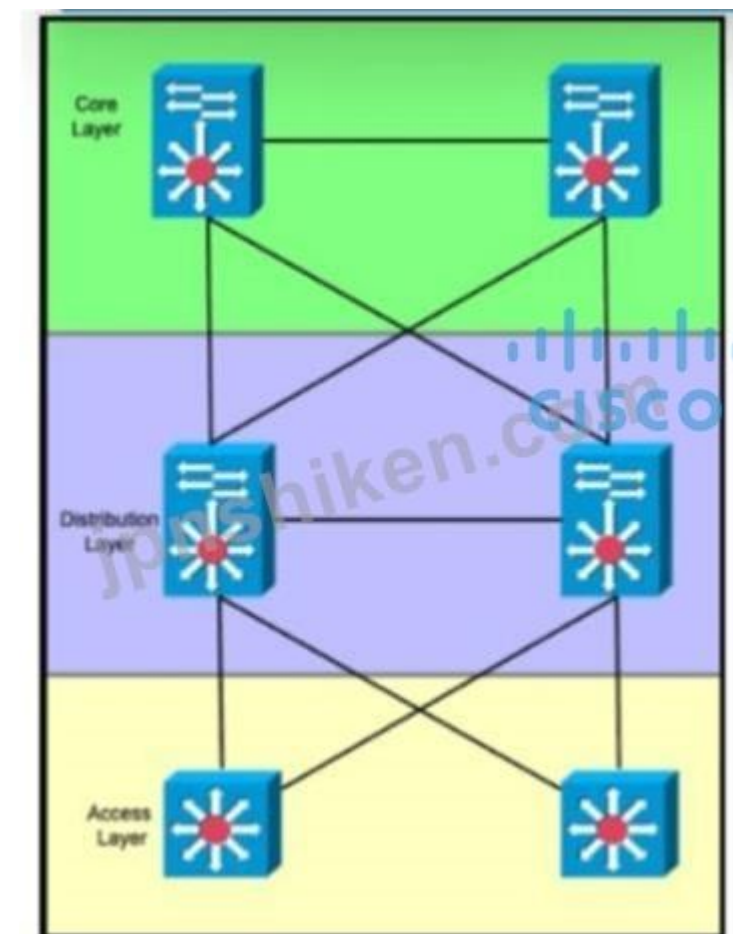
- A. アンダーレイ、MPLS、およびオーバーレイ
- B. プライマリ、バックアップ、負荷分散
- C. 接頭辞、TLOC、およびサービス
- D. インターネット、MPLS、バックアップ

正解: ([正解を表示します](#))

Cisco SD-WAN uses the Overlay Management Protocol to exchange control-plane information between WAN Edge routers and vSmart controllers. OMP advertises three major categories of routing information:

OMP routes, TLOC routes, and service routes. OMP routes represent reachability to prefixes in service VPNs, such as connected, static, OSPF, EIGRP, BGP, or other routes learned at the local site. TLOC routes describe transport locations, which bind a system IP, color, and encapsulation to a WAN transport attachment. Service routes indicate service insertion capabilities, such as firewalls or other network services that can be reached through a particular site. This separation allows Cisco SD-WAN to build a secure overlay over any public or private transport while keeping transport reachability, service-side prefixes, and inserted services distinct. The other answer choices use generic terms such as underlay, MPLS, Internet, backup, or primary routes, but those are not the formal OMP advertisement categories. Therefore, the correct answer is prefix, TLOC, and service.

質問: 65



展示を参照してください。エンジニアは、ルーティング プロトコルとして EIGRP を使用して、マルチキャンパス レイヤー 3 インフラストラクチャを設計しています。設計では、ダウンリンクが発生した場合にクエリにすばやく応答し、不要なクエリを防ぎ、トラフィックがアクセス レイヤーを不必要に通過しないようにする必要があります。エンジニアがネットワーク設計に対して実行する必要がある 2 つのアクションはどれですか。(2 つ選択してください。)

- A. コア層スイッチをスタブ ルータとして設定します。
- B. コア層へのルートを集約するようにディストリビューション層スイッチを設定します。
- C. アクセス レイヤー スwitchをスタブ ルータとして設定します。
- D. アクセス層スイッチとコア層スイッチをスタブ ルータとして設定します。
- E. ディストリビューション レイヤーへのルートを集約するようにアクセス レイヤー スwitchを設定します。

正解: **B,C** ([コメントを發表する](#))

The design should summarize routes from the distribution layer toward the core and configure access-layer switches as EIGRP stubs. EIGRP query scope is a major design consideration in large routed campus networks. When a route is lost, EIGRP may send queries to neighbors to find an alternate path. If the network is not bounded properly, these queries can spread widely, slow convergence, and create stuck-in-active risk.

Configuring the access layer as EIGRP stub tells upstream routers that the access device is not a transit path to other networks, so unnecessary queries are not sent into the access layer. Summarizing at the distribution layer toward the core hides access-link instability and provides a smaller, more stable routing table to the rest of the network. Summarizing at the access layer is usually less useful when many VLANs are aggregated at distribution. Configuring the core as stub is inappropriate because the core should provide transit. Reference topics: EIGRP stub, route summarization, query boundary, campus hierarchy, DUAL convergence.

質問: **66**

クライアントはモデル駆動型テレメトリに移行しており、定期的な更新が必要です。ネットワークアーキテクトは、この設計で何を考慮する必要がありますか？

- A. データ内の変更を含む更新は、変更が発生した場合にのみ送信されます。
- B. 空のデータサブスクリプションは空の更新通知を生成しません。
- C. 定期的な更新には、サブスクライブされているデータの完全なコピーが含まれます。
- D. プライマリプッシュアップデートはすぐに送信され、遅延させることはできません。

正解: ([正解を表示します](#))

For model-driven telemetry with periodic updates, the architect must understand that each periodic update sends the subscribed data at the configured interval. Cisco telemetry guidance distinguishes periodic subscriptions from on-change subscriptions. A periodic subscription streams a full copy of the subscribed data element or table at each interval, even if the underlying state has not changed. That behavior is useful when the collector needs regular time-series samples, trending, or baseline monitoring, but it creates predictable bandwidth and collector-load requirements. On-change subscriptions are different: they publish updates when the subscribed state changes and are more efficient for state changes such as interface status, memory thresholds, or protocol adjacency movement. Empty subscription behavior and initial update timing may matter in specific implementations, but they are not the main design consideration in this question. Since the customer requires periodic updates, the design must size the transport, collector, retention system, and subscription intervals around repeated full data delivery. Reference topics: model-driven telemetry, periodic subscription, on-change subscription, YANG-modeled operational data, collector sizing.

質問: **67**

cisco TelePresence システムをインストールして以来、この会社では、システムの使用中に他のアプリケーションで応答の問題が発生しています。その結果、同社はアーキテクトに QoS ソリューションの推奨を依頼しました。お客様は現在、CBWFQ ポリシーを使用して、速度 100 Mbps のインターネット接続のトラフィックを管理しています。アーキテクトは、リアルタイム トラフィックの完全優先のためにどのリンク容量制限を選択する必要がありますか？

- A. 25Mbps
- B. 50Mbps
- C. 33Mbps
- D. 75Mbps

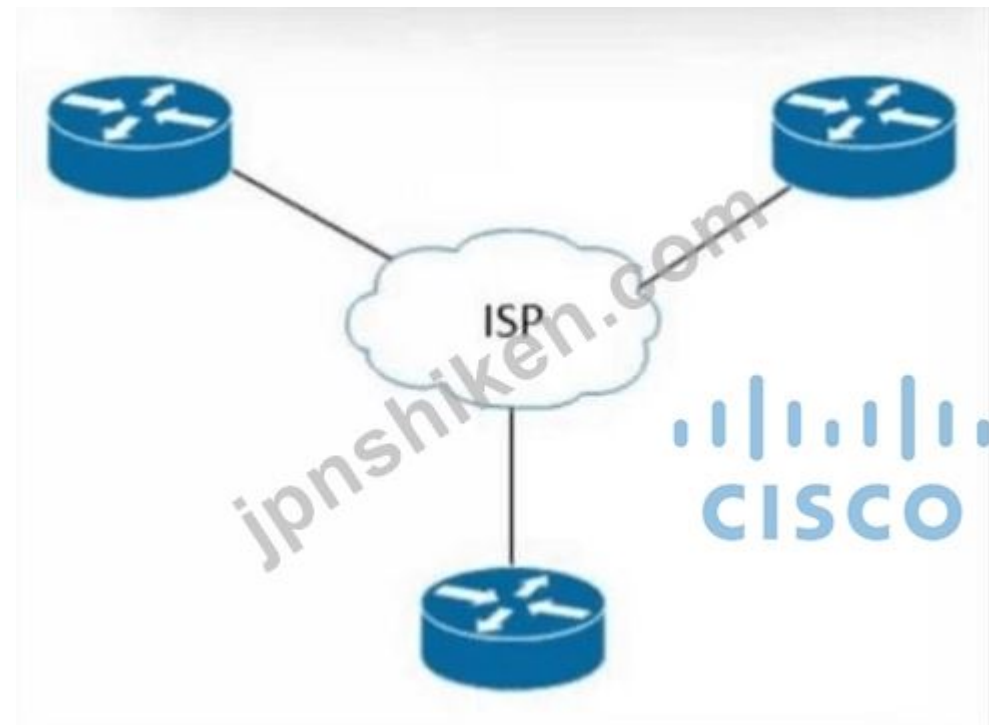
正解: ([正解を表示します](#))

The strict-priority queue should be limited to approximately 33 Mbps on a 100 Mbps Internet connection.

Cisco QoS design guidance warns that strict priority traffic must be bounded, because excessive priority traffic can starve other queues and cause poor application performance for non-real-time traffic. For voice and interactive video, priority queuing is appropriate, but the priority class should normally be constrained to about one-third of the link capacity in a converged enterprise design. That makes 33 Mbps the best choice for a 100 Mbps link. A 75 Mbps strict-priority allocation is far too high and would leave inadequate bandwidth for routing, signaling, business data, and default traffic. A 50 Mbps allocation still gives real-time media excessive control of the interface. A 25 Mbps allocation may work in some environments but is not the standard upper-limit design answer for strict priority in this Cisco context. The architect should also classify TelePresence accurately, mark traffic consistently, and police the priority queue to protect the WAN.

Reference topics: LLQ, CBWFQ, strict priority limits, TelePresence QoS, WAN congestion management.

質問: 68



図を参照してください。予算の制約により、顧客はデバイスごとに1つのLANインターフェイスと1つのWANインターフェイスを備えたWANルーターを購入することにしました。追加のWANリンクを購入せずに高可用性を確保するには、3つのサイトを接続する必要があります。顧客はどの設計展開を選択する必要がありますか？

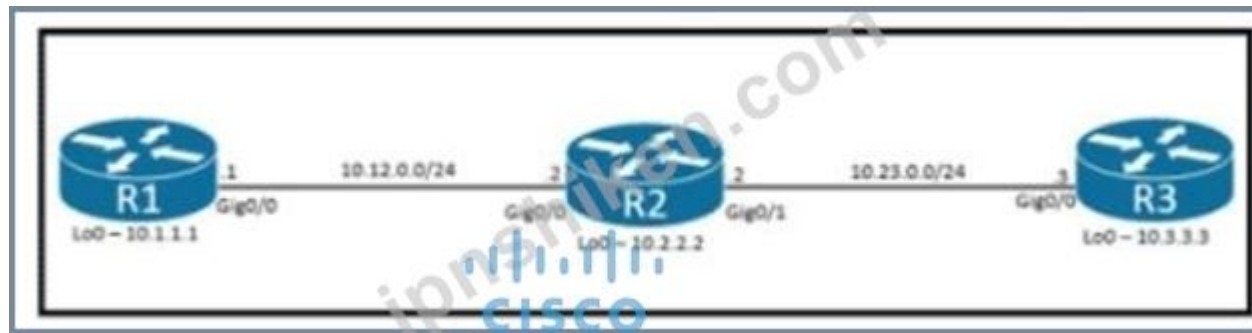
- A. シングルホームフルメッシュ
- B. シングルホームのハブアンドスポーク
- C. デュアルホームハブアンドスポーク
- D. デュアルホームフルメッシュ

正解: ([正解を表示します](#))

A single-homed full-mesh design is the best match when each WAN router has only one WAN interface, no additional WAN links can be purchased, and the customer still wants resilient site-to-site reachability. Single-homed means each site has one physical WAN attachment. Full mesh describes the logical connectivity model between sites, where each site can communicate directly with the others across the WAN service or overlay rather than relying completely on a central hub for intersite traffic. A hub-and-spoke design is simpler, but it creates dependency on the hub for spoke-to-spoke communication and provides poorer resiliency for a three-site design. Dual-homed hub-and-spoke and dual-homed full mesh both require additional WAN connectivity or interfaces, which violates the stated budget and hardware constraints. The best available design is therefore to use the single WAN attachment at each site while building a full-mesh logical topology over the provider or overlay service.

Reference topics: WAN topology design, single-homed connectivity, full mesh, hub-and-spoke limitations, high availability under interface constraints.

質問: 69



図を参照してください。顧客は、OSPF を実行している R1 と R3 間のデータ プレーンの最大稼働時間を要求しています。R2 のルーティング プロセスにメンテナンスが必要な場合、高可用性のために設計にどのソリューションを含める必要がありますか。

- A. すべてのルータ上のBFD
- B. R1 と R3 でのノンストップ転送
- C. R3のみのノンストップ転送
- D. すべてのルータでグレースフルリスタート

正解: [\(正解を表示します\)](#)

Graceful restart is the correct high-availability design when the routing process on an intermediate OSPF router must undergo maintenance while the data plane remains available. Cisco describes graceful restart as a mechanism that allows a restarting router to continue forwarding traffic while neighbors temporarily preserve adjacency and route information. In OSPF, helper-capable neighbors assist the restarting router by maintaining forwarding state during the restart interval, preventing unnecessary route withdrawal and reconvergence. BFD is a failure-detection mechanism, so it would actually accelerate detection of a failure rather than preserve the forwarding path during planned process maintenance. Nonstop forwarding must be enabled on the device performing the control-plane restart and supported by neighbors, but the broad answer choice that captures the required routing-protocol behavior across the design is graceful restart on the participating routers. Configuring it only on R1 and R3 or only on R3 would not support the R2 process maintenance scenario. Reference topics: OSPF graceful restart, NSF, helper mode, control-plane maintenance, data-plane availability.

質問: 70

帯域外管理ネットワークを設計する際に従う必要がある2つのベストプラクティスはどれですか。 2つ選択してください。)

- A. アクセス制御を適用します
- B. ネットワーク統合を促進する
- C. 管理ネットワークを使用してデータをバックアップします
- D. 管理ネットワークがデータネットワークのバックアップであることを確認します
- E. ネットワークの分離を確保

正解: A,E ([コメントを發表する](#))

An out-of-band management network should be isolated from the production data network and protected with strict access control. Cisco SAFE design principles treat the management plane as a sensitive function because it provides administrative access to infrastructure devices. Isolation ensures that production outages, routing problems, or compromised user segments do not automatically remove the operator's ability to manage routers, switches, firewalls, and controllers. Access control ensures that only authorized administrators and management systems can reach the management interfaces and protocols. The management network should not be treated as a backup data path, because that undermines separation and can expose management devices to production traffic risks. It should also not be used as a general-purpose data backup network. "Facilitate network integration" is too vague and can conflict with the isolation requirement. Therefore, the best practices are to enforce access control and ensure network isolation. In a mature design, this normally includes dedicated management interfaces or VRFs, jump hosts, AAA, logging, and tightly controlled management protocol access.

質問: 71

エンジニアは、サービスプロバイダーとのデュアルBGPピアリングソリューションの設計を担当しています。設計は次の条件を満たす必要があります。

*ルーターは、サブネットマスクが/ 24より大きいプレフィックスを学習しません。

*ルーターは、マスクの長さのみに基づいて、ルーティングテーブルに含めるルートを決めます。

*ルーターは、サービスプロバイダーの構成に関係なく、この選択を行います。
エンジニアはどのソリューションを設計に含める必要がありますか？

- A. ルートマップとアクセスリストを使用して目的のネットワークをブロックし、ルートマップをインバウンドのBGPネイバーに適用します。
- B. ルートマップとプレフィックスリストを使用して目的のネットワークをブロックし、ルートマップをBGPネイバーのアウトバウンドに適用します。
- C. IPプレフィックスリストを使用して目的のネットワークをブロックし、IPプレフィックスリストをBGPネイバーのアウトバウンドに適用します。
- D. IPプレフィックスリストを使用して目的のネットワークをブロックし、IPプレフィックスリストをインバウンドのBGPネイバーに適用します。

正解: [\(正解を表示します\)](#)

An inbound IP prefix list is the correct tool because the routers must decide which prefixes to accept from the service provider based only on prefix length. Prefix lists are designed to match network prefixes and mask-length ranges using operators such as le and ge. Applying the prefix list inbound to the BGP neighbors prevents routes with masks longer than /24 from entering the local BGP table, regardless of what the service provider advertises. This conserves memory and ensures the customer routing policy is enforced locally. An access list is a poor fit because it does not natively express prefix-length logic as cleanly as a prefix list.

Applying the policy outbound would filter routes the customer sends to the provider, which is the wrong direction; the requirement is about routes the customer learns. A route map could reference a prefix list, but the simplest and most exact answer is to use an IP prefix list inbound. Therefore, the design should apply an inbound prefix list to both BGP peers to block prefixes greater than /24. Reference topics: BGP route filtering, prefix lists, inbound policy, maximum accepted prefix length.

質問: 72

アーキテクトは、スパニングツリープロトコルを利用してループのないトポロジを確保するネットワークを設計しています。ネットワークは、エンドユーザーがテスト目的で独自のネットワークスイッチを接続する必要があるエンジニアリング環境をサポートします。スパニングツリートポロジがこれらの不正なスイッチの影響を受けないようにするために、アーキテクトはどの機能を設計に含める必要がありますか？

- A. BPDUスキュー検出
- B. BPDUガード
- C. ループガード
- D. ルートガード

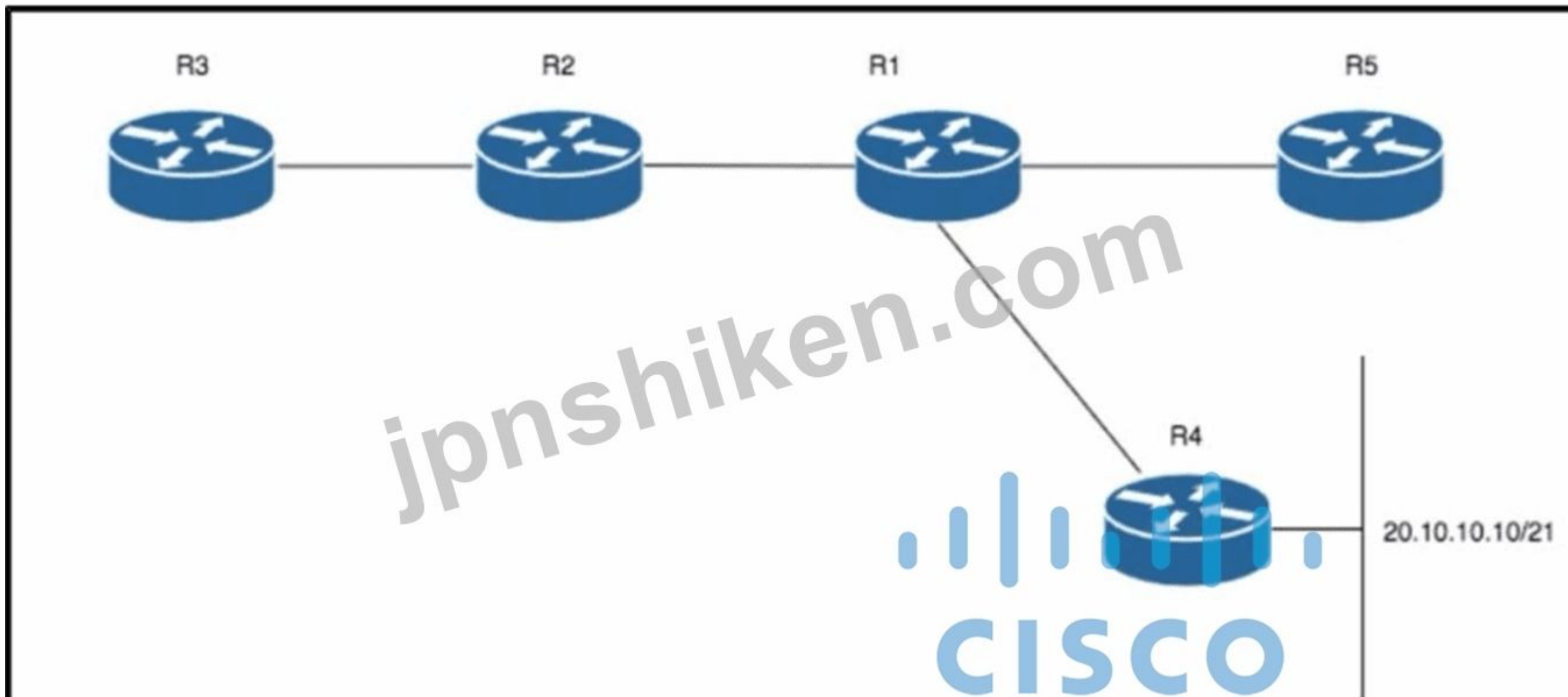
正解: **D** ([コメントを發表する](#))

Root guard is the correct feature when the design must allow switches to be attached for testing but prevent those switches from influencing the spanning-tree root placement. Cisco root guard protects the intended STP root by placing a port into a root-inconsistent state if superior BPDUs are received on a port where the root should never appear. This directly addresses the risk created by rogue or lab switches connected by end users:

if a user switch advertises a better bridge priority, the production access topology is not allowed to reconverge around that unauthorized device. BPDU guard is usually used on PortFast edge ports to err-disable a port that receives any BPDU, which is excellent for strict host-only ports but less aligned with the wording that users may connect their own switches for testing. Loop guard protects root or alternate ports from losing BPDUs and accidentally forwarding. BPDU skew detection is a diagnostic feature, not the correct design control.

Therefore, root guard is the appropriate protection against rogue switches affecting STP topology. Reference topics: STP root placement, root guard, superior BPDUs, campus Layer 2 protection.

質問: 73



図を参照してください。ネットワーク設計者がEIGRPに基づいたネットワーク設計を準備しています。ルータはCat6aケーブルで接続されており、距離の関係でルータ間の接続速度は10Mbpsに制限されています。パイロットフェーズ中にDUAL-3-SIAエラーメッセージが表示されました。安定した設計を作成するために、エンジニアはどのような対策を講じる必要がありますか？

- A. R4で毒逆転を有効にする。
- B. R2上にサマリー ルートを作成します。
- C. R1のスプリット ホライズンを無効にします。
- D. R4上にSTUB領域を設定します。

正解: ([正解を表示します](#))

The correct action is to configure EIGRP stub behavior on R4. The DUAL-3-SIA message indicates a Stuck-in-Active condition, which occurs when EIGRP sends a query and does not receive a reply within the expected time. Cisco explains that EIGRP stub routing is used to improve stability and reduce resource usage by preventing unnecessary queries from being sent to remote or edge routers. In the described topology, R4 is connected over a limited 10 Mbps inter-router link and should not be treated as a transit router for the rest of the routing domain. Making R4 an EIGRP stub limits the query scope and reduces the chance of active routes waiting on a slow or constrained neighbor. Poison reverse and split-horizon changes address loop prevention and hub-and-spoke update behavior, but they do not directly solve the SIA instability. A summary on R2 may help scaling but does not identify the edge-router query boundary as precisely as EIGRP stub. Reference topics: EIGRP DUAL, Stuck-in-Active, EIGRP stub routing, query scoping, WAN edge stability.



図を参照してください。ある建築家がレイヤー3キャンパスネットワークを設計しています。設計では、ネットワークの不安定性を隠蔽し、ネットワークのオーバーヘッドを削減し、重要なデバイスのメモリを節約する必要があります。建築家はどの経路集約ソリューションを選択すべきでしょうか？

- A. アグリゲーション層は、アクセス層へのデフォルトルートを実バタイズする必要があります。VLANサブネットは、アグリゲーション層で10.0.0.0/16に集約され、コア層に実バタイズされる必要があります。
- B. コアレイヤは、集約レイヤへのデフォルトルートを実バタイズする必要があります。VLANサブネットは、アクセスレイヤで10.0.0.0/16に集約され、集約レイヤに実バタイズされる必要があります。
- C. アグリゲーションレイヤは、コアレイヤへのデフォルトルートを実バタイズする必要があります。VLANサブネットは、アグリゲーションレイヤで10.0.0.0/16に集約され、アクセスレイヤに実バタイズされる必要があります。
- D. コアレイヤは、集約レイヤへのデフォルトルートを実バタイズする必要があります。VLANサブネットは、集約レイヤで10.0.0.0/16に集約され、コアレイヤに実バタイズされる必要があります。

正解: [\(正解を表示します\)](#)

The correct design is to advertise a default route from the aggregation layer toward the access layer and summarize the VLAN subnets at the aggregation layer toward the core. Cisco hierarchical campus design places summarization at aggregation or distribution boundaries because that is where access-layer routes are collected before entering the core. Summarizing the VLAN ranges into 10.0.0.0/16 hides individual access VLAN instability from the core, reducing routing table entries, protocol churn, and memory consumption on core devices. Sending a default route from aggregation to access also keeps access switches simple and prevents them from carrying unnecessary campus routing detail. This design contains failures and convergence events inside the appropriate block. Summarizing at the access layer is less practical when multiple VLANs aggregate at distribution. Advertising a default toward the core is backwards, because the core must know the summarized access block. Reference topics: hierarchical campus routing, aggregation- layer summarization, default routing to access, route-table reduction, convergence containment.

質問: 75

アーキテクトは、100台を超えるルーターを含むネットワーク用のマルチキャストソリューションを設計しています。アーキテクトは、複数のマルチキャストドメインを作成し、ネットワーク内でPIM-SMトラフィックのバランスを取ることを計画しています。アーキテクトはどのテクノロジーを設計に含める必要がありますか？

- A. DVMRP
- B. IGMP
- C. MOSPF
- D. MSDP

正解: **D** ([コメントを发表する](#))

MSDP is the correct technology when a large multicast design uses multiple PIM Sparse Mode domains and needs to exchange active source information between those domains. Cisco describes Multicast Source Discovery Protocol as the mechanism used by RPs in different PIM-SM domains to learn about multicast sources outside their own domain. When a source becomes active, the local RP can advertise Source-Active information to MSDP peers, allowing receivers in other domains to discover and join that source. This supports domain separation and can help balance multicast control-plane scope in large networks. DVMRP and MOSPF are legacy multicast routing approaches and are not the design mechanism for modern PIM-SM domain interconnection. IGMP is used between hosts and local multicast routers for receiver membership signaling; it does not exchange source information between multicast domains. Therefore, a design with more than 100 routers, several multicast domains, and PIM-SM traffic balancing should include MSDP. The architect should also design RP placement, Anycast RP behavior, peer mesh scaling, and source filtering.

Reference topics: MSDP, PIM-SM domains, Source-Active messages, interdomain multicast.

質問: 76

カスタマーソリューションでは、WANを介したストリーミングマルチメディアをサポートするためにQoSが必要です。アーキテクトは、ホップごとの動作を使用することを選択します。エンジニアは、ランチャサイト間を移動するトラフィックをマークするためにどのソリューションを使用する必要がありますか？

- A. DSCP EFを使用したLLQ
- B. DSCP AF3を使用したCBWFQ
- C. DSCP AF2を使用したCBWFQ
- D. DSCP AF4を使用したLLQ

正解: ([正解を表示します](#))

For streaming multimedia over a WAN using DiffServ per-hop behavior, the best answer is CBWFQ with DSCP AF3. Cisco QoS design separates real-time interactive voice from streaming multimedia.

Voice bearer traffic commonly uses EF and a strict priority queue because it is extremely sensitive to delay and jitter.

Streaming multimedia, however, is usually buffered and bandwidth-intensive; it should receive assured bandwidth and controlled drop behavior rather than strict priority treatment. The Assured Forwarding classes are designed for this kind of preferential forwarding within DiffServ. AF3 is commonly associated with multimedia streaming in Cisco QoS baseline models, while CBWFQ provides bandwidth guarantees without allowing the class to starve other traffic. LLQ with EF is too aggressive for streaming multimedia and should be reserved for voice bearer traffic or similarly strict real-time flows. AF2 is more commonly used for transactional data or lower-priority assured service, depending on the model. The correct marking and queuing design is therefore CBWFQ with DSCP AF3. Reference topics: DiffServ, per-hop behavior, CBWFQ, DSCP AF classes, multimedia streaming QoS.

有効的な**300-420J**問題集はJPNTTest.com提供され、**300-420J**試験に合格することに役に立ちます！JPNTTest.comは今最新**300-420J**試験問題集を提供します。JPNTTest.com 300-420J試験問題集はもう更新されました。ここで**300-420J**問題集のテストエンジンを手に入れます。最新版のアクセス、<https://www.jpntest.com/shiken/300-420J-mondaishu> **457**問、**30%ディスカウント**、特別な割引コード：**JPNshiken**」

質問: 77

independent of underlying platform

platform dependent

standards dependent

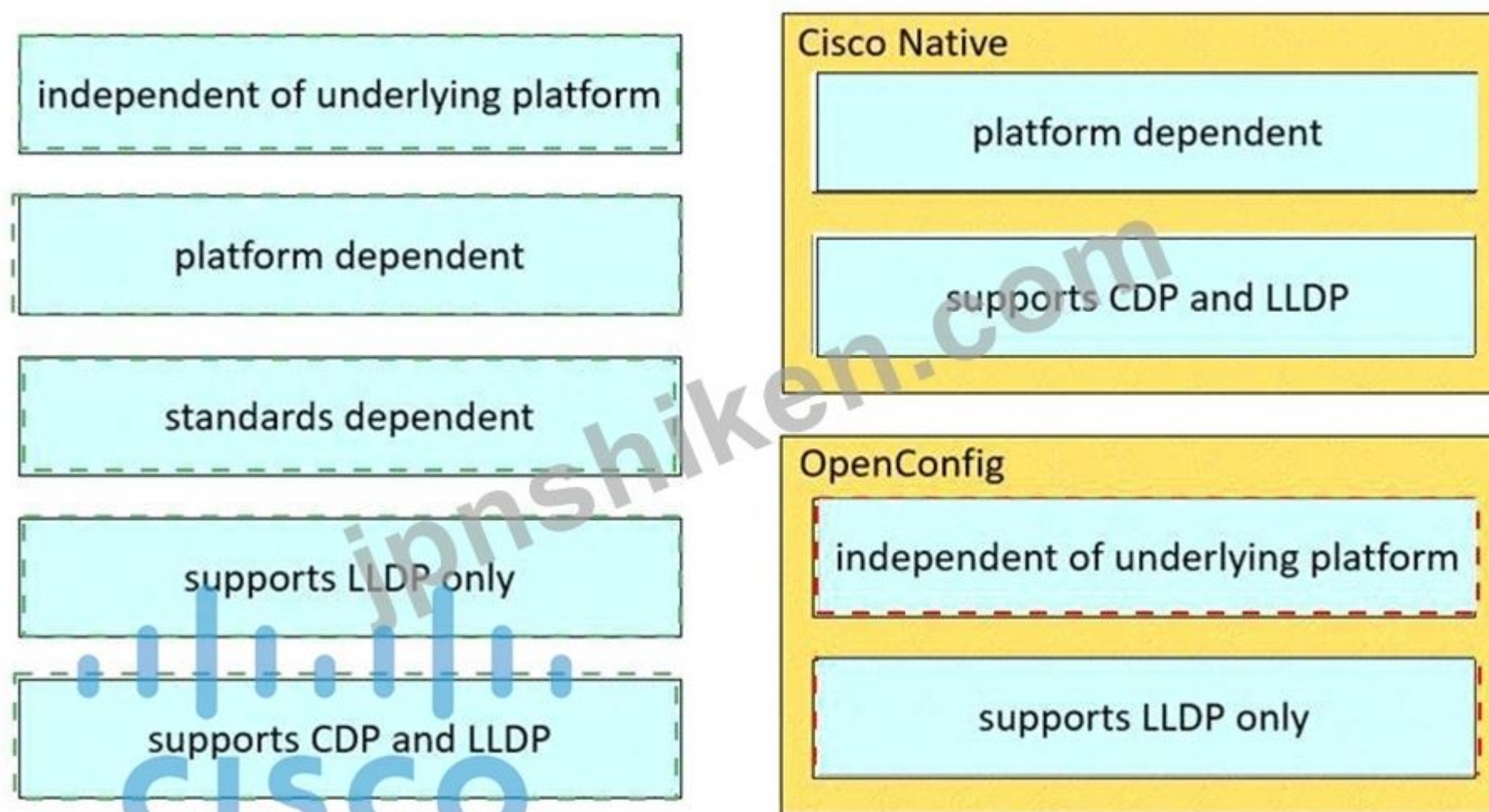
supports LLDP only

supports CDP and LLDP

Cisco Native

OpenConfig

正解:



Explanation:

This YANG mapping is based on how IETF, OpenConfig, and Cisco native models differ in portability and schema organization. IETF models are produced through the standards process and are appropriate when the same management structure is expected across vendors for common protocols. OpenConfig models are also vendor-neutral, but they are built around a consistent operational model used heavily for streaming telemetry and multivendor automation. Cisco native models mirror Cisco platform configuration more closely and expose device-specific capabilities that may not exist in common models. The selected answers keep those boundaries clear. This matters because a NETCONF or RESTCONF payload is validated against the YANG schema; a payload written for one model family cannot be assumed to work against another. In design terms, the model family should be selected after confirming device software release support, feature coverage, operational-state requirements, and vendor diversity. Open models give portability; native models give feature completeness. Reference topics: IETF YANG, OpenConfig YANG, Cisco native YANG, schema validation, automation model selection.

質問: 78

ネットワーク設計者はLANでテレビサービスを有効にしています。ソースは

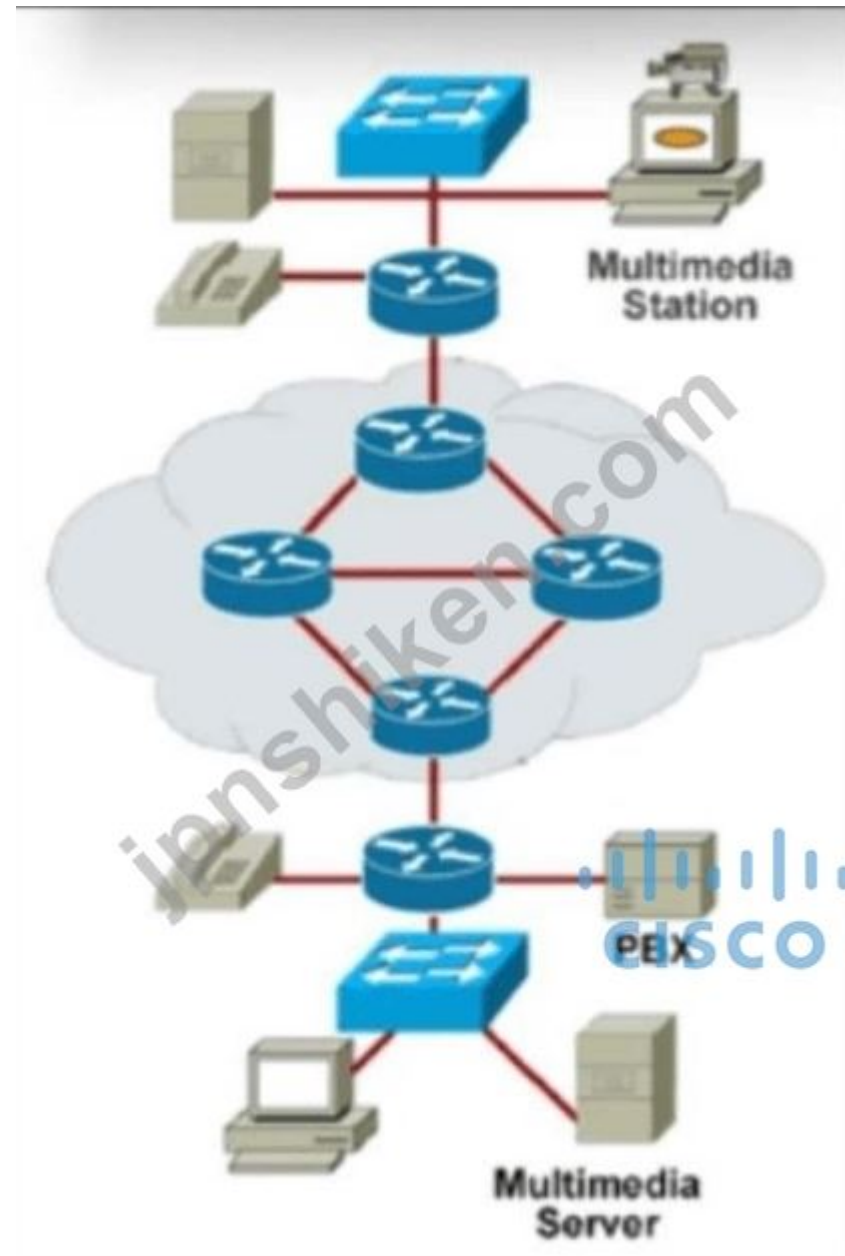
239.1.1.1 グループ IP アドレス。ネットワークでは高密度モードは許可されていません。マルチキャストは、LAN セグメント内のすべてのネットワーク デバイスですでに有効になっています。設計を完了するには、アーキテクトはどのようなアクションを実行する必要がありますか？

- A. PIM SSM を有効にします。
- B. PIM Auto-RP を有効にします。
- C. PIM エニーキャスト RP を有効にします。
- D. PIM BSR を有効にします。

正解: ([正解を表示します](#))

PIM BSR is the best final design action because the source uses group 239.1.1.1, which is an administratively scoped ASM multicast group, and dense mode is not allowed. In PIM sparse mode, ASM groups require a rendezvous point so receivers can join the shared tree and sources can register. PIM SSM is not the right match for 239.1.1.1 because SSM normally uses the 232.0.0.0/8 group range and requires receivers to know the source address. Auto-RP can distribute RP information, but traditional Auto-RP discovery relies on multicast groups that often require sparse-dense behavior, which conflicts with the requirement that dense mode is not allowed. BSR is the standards-based PIM mechanism for dynamic RP discovery in sparse-mode networks and avoids static RP configuration on every router. Anycast RP is an RP redundancy design but still requires RP discovery or static RP configuration. Reference topics: PIM sparse mode, Bootstrap Router, ASM, RP discovery, administratively scoped multicast.

質問: 79



展示を参照してください。エンジニアは、QoS 設計によってアプリケーションの帯域幅が保証され、アプリケーションが遅延要件をサポートするために特定のタイプのサービスを要求できることを確認する必要があります。エンジニアはどのソリューションを選択する必要がありますか？

- A. RSVP 付き Diffserv
- B. RSVP 付き IntServ
- C. DSCP を使用した Diffserv
- D. DSCP 付き IntServ

正解: ([正解を表示します](#))

IntServ with RSVP is the correct model when an application must request a specific service level and receive guaranteed bandwidth or delay treatment. Cisco describes the Integrated Services model as a flow-based QoS architecture in which applications signal their resource requirements to the network. RSVP is the signaling protocol used to reserve resources along the path. This differs from DiffServ, which marks packets and applies per-hop behavior without creating per-flow reservations. The question states that the application can request a particular type of service to support its delay requirements and that bandwidth must be guaranteed.

Those are classic IntServ and RSVP characteristics. DiffServ with DSCP is scalable and widely used in enterprise networks, but it cannot guarantee resources end to end unless every hop is engineered and provisioned accordingly. IntServ with DSCP is not the standard model because RSVP is the reservation mechanism. Reference topics: Integrated Services, RSVP, resource reservation, guaranteed service, QoS architecture, delay-sensitive applications.

質問: 80

ポリシーが定義されていない場合、Cisco Catalyst SD-WANアーキテクチャにおいてOMPはどのように動作しますか？

- A. WANエッジルーターが中央拠点を通して通信するためのハブアンドスポークトポロジを可能にする
- B. WANエッジルーターが中央拠点から遠隔拠点へ通信するためのポイントツーポイントトポロジを可能にするため
- C. すべてのWANエッジルーターがフルメッシュトポロジを使用して通信できるようにする
- D. WANエッジルーター間のすべての通信をブロックする

正解: ([正解を表示します](#))

Without centralized policy, OMP allows all WAN Edge routers in the same VPN to communicate in a full- mesh overlay. Cisco Catalyst SD-WAN uses OMP as the control-plane protocol between WAN Edge routers and SD-WAN controllers. WAN Edge devices advertise OMP routes, TLOC routes, and service routes to the controllers, and the controllers distribute the reachable routes back to the relevant edge devices. Cisco policy documentation describes the default unpoliced behavior as permissive: routes are accepted and advertised so sites in the same VPN learn reachability to each other. That default route distribution produces a full-mesh data-plane capability, even though control connections remain hubbed through the controllers. Hub-and- spoke, regional isolation, point-to-point restriction, or service-chained topologies require centralized control or data policy to filter routes or steer traffic. Blocking all communication is the opposite of the default behavior. Reference topics: Cisco SD-WAN OMP, default route distribution, full-mesh overlay, centralized control policy, VPN segmentation.

質問: 81

ある企業は、複数の部門の移転を支援するため、本社を拡張しています。ネットワークはこれまでRIPで運用されていましたが、企業の成長に伴い、エンジニアリングチームはより堅牢なルーティングプロトコルをサポートする必要があると判断しました。チームは、以下の要件を満たすルーティングソリューションの設計に取り組んでいます。

- * ネットワーク分離をサポートし、セグメント間で要約を行います
- * 高速な収束を保証します
- * 拡張性を提供する

ネットワークチームはどの設計を採用すべきか？

- A. スタブエリアを使用して EIGRP をデプロイします。
- B. 複数のエリアで OSPF を展開します。
- C. 少なくとも 2 つのセグメントを接続する仮想リンクを使用して OSPF をデプロイします。
- D. 変更されたK値を使用してEIGRPを展開します。

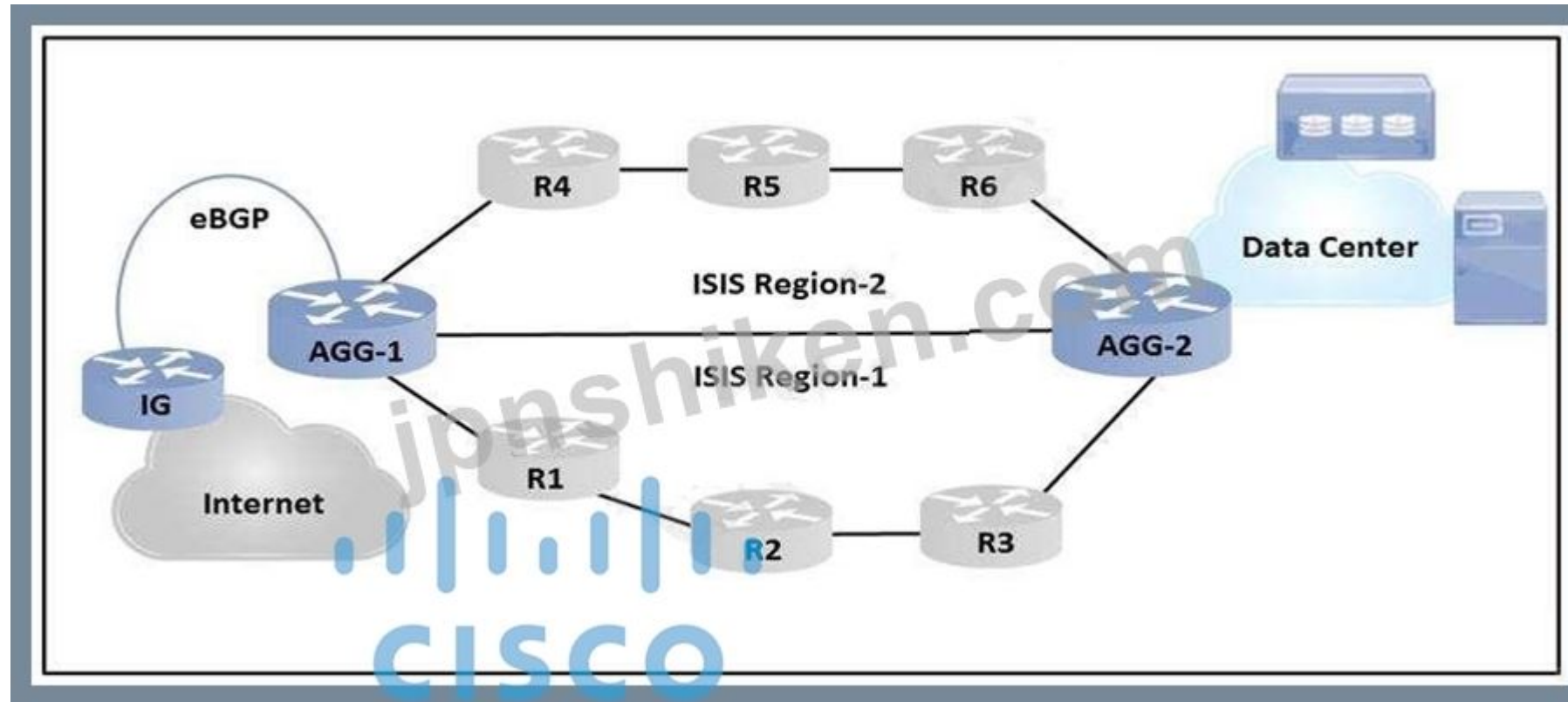
正解: ([正解を表示します](#))

The correct routing design is OSPF with multiple areas. The company is moving away from RIP because the network has grown and now requires segmentation, summarization, fast convergence, and scalability. Cisco positions OSPF as a link-state IGP suitable for larger enterprise networks, and OSPF areas are the standard mechanism for containing topology information and scaling the link-state database. Area Border Routers can summarize routes between areas, reducing routing table size and limiting the scope of link-state flooding and SPF recalculation. This directly satisfies the requirement to segregate departments while summarizing between segments. EIGRP can converge quickly and summarize, but "stub areas" are not an EIGRP construct; EIGRP uses stub routers. OSPF virtual links are

not a primary segmentation design tool; they are used to repair backbone connectivity issues. Modifying EIGRP K values changes metric calculation but does not solve segmentation or summarization. Reference topics: OSPF multi-area design, ABR summarization, SPF scope control, enterprise IGP scalability.

質問: 82

展示品を参照してください。



アーキテクトは、エンタープライズ顧客向けの IGP ソリューションを設計する必要があります。設計では、次の内容をサポートする必要があります。物理リンクフラップの影響は最小限に抑えられます。

アクセス ルータは、リンク障害発生後すぐに収束する必要があります。

建築家が設計に含めるべき IS-IS ソリューションはどれですか? (2 つ選択してください。)

- A. BGP から IS-IS への再配布を使用して、レベル 1 エリア内のすべてのインターネット ルートをアドバタイズします。
- B. IS-IS インターフェイスとループバック IP アドレスをインターネットとデータセンターに向けてアドバタイズします。
- C. SPF および PRC 間隔を短縮して、収束時間を改善します。
- D. ネットワーク全体でレベル 1/レベル 2 隣接関係を確立するように、すべてのアクセス ルータと集約ルータを設定します。
- E. アクセス ルータを設定してレベル 1 隣接関係を確立し、集約ルータを設定してレベル 1/レベル 2 隣接関係を確立します。

正解: ([正解を表示します](#))

The correct IS-IS design combines tuned SPF/PRC behavior with a proper Level 1 and Level 2 hierarchy.

Reducing SPF and PRC intervals, within platform and stability limits, helps routers react more quickly to topology changes after a link failure. However, fast timers alone are not enough; the hierarchy must contain instability so physical link flaps in the access layer have limited impact on the rest of the network. Access routers should form Level 1 adjacencies inside their local area, while aggregation routers should operate as Level 1/Level 2 routers to connect the access area to the Level 2 backbone. This design gives access routers a simple local view and uses aggregation as the interarea boundary.

Running all access and aggregation routers as L1/L2 across the network increases the size of the backbone and exposes more routers to wider flooding, which is the opposite of the requirement.

Redistributing Internet routes into Level 1 is also poor design because it bloats the access area. Therefore, the right choices are C and E. Reference topics: IS-IS hierarchy, L1/L2 design, SPF and PRC tuning, access-layer convergence, fault containment.

質問: 83

どのタイプのランデブーポイント展開が標準ベースであり、動的 RP 検出をサポートしていますか？

- A. ブートストラップルータ
- B. エニーキャストRP
- C. 自動RP
- D. 静的RP

正解: **A** ([コメントを发表する](#))

Bootstrap Router is the standards-based mechanism that supports dynamic RP discovery in a PIM sparse-mode domain. Cisco documentation distinguishes Auto-RP, which is Cisco-proprietary, from BSR, which is defined for standards-based RP distribution. With BSR, candidate RPs advertise their availability, and the elected bootstrap router distributes RP-set information throughout the PIM domain. This removes the operational burden of manually configuring the same static RP information on every multicast router. Anycast-RP provides RP redundancy and load sharing but does not by itself provide standards-based dynamic RP discovery across the domain. Static RP is simple but not dynamic. Auto-RP can dynamically distribute RP information, but it is not the standards-based answer. The design value of BSR is most visible in large multicast domains, where routers need consistent RP information and manual changes are risky. Therefore, when the question asks for a standards-based RP deployment that supports dynamic RP discovery, the correct answer is bootstrap router. Reference topics: PIM sparse mode, Bootstrap Router, candidate RP, RP-set distribution, dynamic RP discovery.

質問: **84**

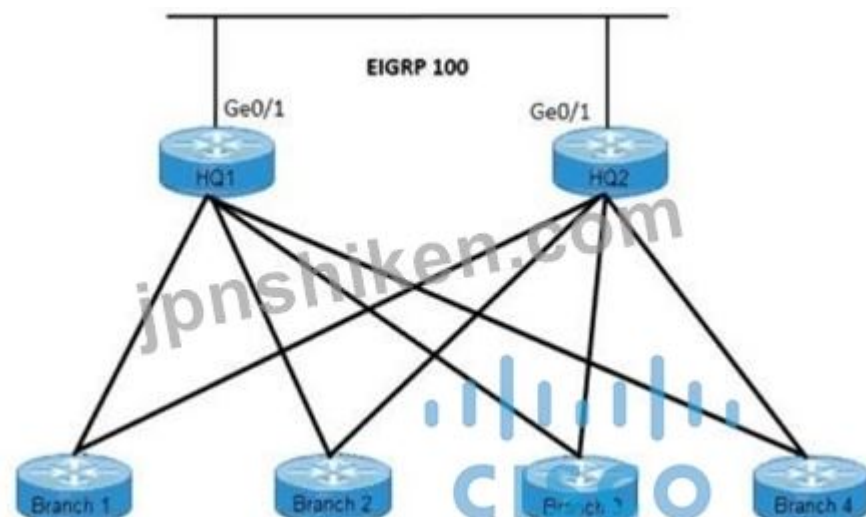
Cisco SO-Access アーキテクチャ内で VN 内トラフィックのフィルタリングと制御の機能を提供する機能はどれですか？

- A. スケーラブルなグループ
- B. MAC ACL
- C. プレフィックスリスト
- D. サービスポリシー

正解: **A** ([コメントを发表する](#))

Scalable groups provide intra-VN traffic filtering and control in Cisco SD-Access. Within a virtual network, macro segmentation separates routing tables, but endpoints inside the same VN may still need policy control based on identity, role, or security posture. Cisco SD-Access uses scalable group tags and scalable group ACLs to implement group-based policy. This allows policy to be expressed as user or device group relationships instead of only subnet-based ACL entries. The result is microsegmentation inside the virtual network, where traffic between groups can be permitted or denied consistently across the fabric. MAC ACLs are limited Layer 2 filters and do not provide the identity-based policy model required for SD-Access. Prefix lists are route-filtering tools, not intra-VN access-control mechanisms. Service policies are used for QoS or traffic treatment and do not define group-based segmentation. Scalable groups are therefore the correct policy construct for intra-VN filtering in the fabric. Reference topics: Cisco SD-Access policy, scalable groups, SGT, SGACL, microsegmentation, intra-VN control.

質問: **85**



展示を参照してください。アーキテクトは、顧客向けに安定したスケーラブルな EIGRP ソリューションを作成する必要があります。設計では次のことを行う必要があります。

*帯域幅、メモリ、CPU処理を節約

*最適でないルーティングを防ぐ

*不必要な問い合わせを避ける

建築家はどの2つのソリューションを選択する必要がありますか? (2つ選択してください。)

- A. ルート集約
- B. プレフィックスリスト
- C. リストを配布する
- D. スタブルーティング
- E. 静的再配布

正解: ([正解を表示します](#))

Route summarization and EIGRP stub routing are the two correct design techniques for a stable and scalable EIGRP deployment. Summarization reduces the number of prefixes advertised upstream, conserves bandwidth and memory, hides route flaps behind the summary boundary, and reduces the size of the EIGRP topology table. Stub routing prevents branch or access devices from being used as transit routers and limits the scope of EIGRP queries. This is critical because uncontrolled EIGRP query propagation can increase convergence time and processing load, especially in larger hierarchical networks. Prefix lists and distribute lists can filter routes, but filtering alone does not provide the same query-boundary and topology-hiding benefits as summarization and stubs. Static redistribution increases complexity and can create loop or policy risks if not controlled with tags and metrics. The question specifically asks to conserve bandwidth, memory, and CPU, prevent suboptimal routing, and avoid unnecessary queries; summarization and stub routing are the direct Cisco design answers. The architect should summarize at hierarchy boundaries and configure access or branch routers as EIGRP stubs where they should not provide transit. Reference topics: EIGRP summarization, EIGRP stub routing, query scope, hierarchical routing.

質問: 86

OSPFネットワークでの収束を改善するために実行できる2つのステップはどれですか。 (2つ選択してください。)

- A. 双方向転送検出を使用します
- B. すべてのエリアを1つのバックボーンエリアにマージします
- C. OSPFパラメーターの調整
- D. すべての非バックボーンエリアをスタブエリアにする
- E. 複数のエリアにわたって同じIPネットワークにまたがります。

正解: ([正解を表示します](#))

The two practical OSPF convergence improvements are enabling Bidirectional Forwarding Detection and tuning OSPF timers and throttling behavior. Cisco describes BFD as a protocol designed to provide fast forwarding-path failure detection across media, encapsulations, topologies, and routing protocols. With OSPF, BFD allows a neighbor failure to be detected rapidly without depending only on the OSPF dead interval.

OSPF parameters can also be tuned, including SPF scheduling, LSA generation, and LSA arrival throttling, so the control plane reacts quickly but not recklessly during instability. Merging all areas into one backbone area is not a convergence improvement; it usually increases the size of the link-state database and the impact of topology changes. Making every nonbackbone area stub may reduce external LSA scope, but it is not universally valid and does not directly solve failure detection. Spanning the same IP network across multiple OSPF areas is a poor design and can create operational ambiguity.

Therefore, BFD plus careful OSPF parameter tuning are the appropriate convergence-focused design actions.

質問: 87



図を参照してください。顧客は、単一のCEルータから2つのサービスプロバイダに対して2つのeBGPピアリングを確立しています。顧客は、特定のトラフィックが確実に顧客側に入るようにするためのソリューションを設計するために、アーキテクトを雇いました。

インターフェース gig0/0 を介して s ネットワークに接続します。設計者はどのソリューションを設計に含める必要がありますか？

- A. 好ましくないサービスプロバイダーに対して、より低いMED値を宣伝する。
- B. 優先サービスプロバイダへのASパスに、追加のASを先頭に追加します。
- C. 集約されたルートをより長いプレフィックスに分割し、優先サービスプロバイダに通知します。
- D. 優先サービスプロバイダパスに高いローカル優先度を設定します。

正解: ([正解を表示します](#))

Advertising more-specific prefixes toward the preferred service provider is the correct way to influence inbound traffic to enter through a chosen CE interface. BGP inbound path control is indirect: the enterprise cannot set local preference inside remote autonomous systems, and MED is honored only under specific provider policies. The most deterministic option available to the customer is to advertise a longer-prefix route to the preferred provider because Internet and BGP route selection generally prefer the longest matching prefix before BGP path attributes are considered. If the aggregate is advertised to both providers but a more-specific route is advertised to the provider connected to gig0/0, return traffic for that more-specific destination is attracted through that interface. Lowering MED toward the less preferred provider is the opposite of the desired effect. AS-path prepending toward the preferred provider would make that path less attractive. Local preference affects outbound path selection inside the local AS, not inbound traffic from external networks. Reference topics: BGP inbound traffic engineering, longest-prefix match, route aggregation, AS-path prepending, MED, local preference.

質問: 88

インフラストラクチャチームは、デバイスの共有メモリ使用率を懸念しているため、デバイスの状態を監視する必要があります。デバイスへの影響を制限し、必要なデータを提供するソリューションはどれですか？

- A. IPFIX
- B. 静的テレメトリー
- C. 変更時のサブスクリプション
- D. 定期的なサブスクリプション

正解: ([正解を表示します](#))

An on-change telemetry subscription is the best solution when the team needs device state data while limiting impact on the device. With periodic telemetry, the device sends subscribed data at a configured interval regardless of whether the value has changed. For state such as shared memory utilization, aggressive periodic polling or streaming can create unnecessary CPU, memory, and bandwidth overhead, especially at scale. On-change telemetry sends an update only when the subscribed data changes, or when a relevant event occurs according to the model and platform behavior. This reduces needless data transmission while still providing timely visibility into meaningful state changes. IPFIX is designed for flow information and traffic accounting rather than detailed device state telemetry. Static telemetry is not the best characterization of the efficient event-driven behavior required here. A periodic subscription can provide the data, but it continuously streams updates and therefore has greater device and collector impact. Cisco model-driven telemetry design favors on-change subscriptions when the operational requirement is state awareness with minimal repeated polling.

質問: 89

複数のネットワークロケーションに同じサブネットが存在できるようにするコントロールプレーンテクノロジーはどれですか。

- A. LISP
- B. VXLAN
- C. FabricPath
- D. ISE mobility services

正解: ([正解を表示します](#))

Locator/ID Separation Protocol is the control-plane technology that allows endpoint identity to be separated from endpoint location. In Cisco SD-Access, LISP is used by the fabric control plane to map endpoint identifiers to routing locators, so an endpoint subnet can remain logically consistent even when endpoint reachability is provided from different network locations. This is why LISP is associated with host mobility and with designs where the same subnet or endpoint identity can exist across multiple places without forcing the physical topology to be a single stretched Layer 2 domain. VXLAN is the data-

plane encapsulation used to carry overlay traffic, but VXLAN itself is not the mapping control plane. FabricPath is a Cisco Layer 2 multipath technology and does not provide LISP-style endpoint-location mapping. ISE mobility services are identity and policy functions, not a routing locator mapping system. Therefore, the correct control-plane technology is LISP. Cisco SD-Access design material describes the control plane as the mapping system that maintains endpoint-to-device relationships.

質問: 90

同じ自律システム内のEIGRPネイバー間のルートをフィルタリングする方法はどれですか。

- A. 配布リスト
- B. ポリシーベースのルーティング
- C. リークマップ
- D. ルートのタグ付け

正解: (正解を表示します)

A distribute list is the correct method to filter routes exchanged between EIGRP neighbors within the same autonomous system. In Cisco EIGRP designs, distribute lists can be applied inbound or outbound under the EIGRP process to control which routes are accepted from or advertised to neighbors. The filter can reference an access list, prefix list, or route map depending on platform and configuration style. This allows the designer to limit route visibility without changing the forwarding behavior of packets that are already routed.

Policy-based routing is a data-plane forwarding feature that changes next-hop selection for matching traffic; it does not filter EIGRP route advertisements between neighbors. Leak maps are used with route summarization to selectively advertise more specific routes through a summary, not as the general route-filtering method between EIGRP neighbors. Route tagging is valuable in redistribution designs to identify route origin and prevent feedback loops, but the tag alone does not filter routes unless a route map or distribute list uses it.

Therefore, the correct route-filtering mechanism between EIGRP neighbors is a distribute list.

質問: 91



展示を参照してください。アーキテクトは、IPv4 から IPv6 に移行する顧客向けに ISIS ネットワークを設計しています。

現在のネットワークでは狭いメトリックが使用されており、今後 2 年以内に IPv6 エリアが 10 に増加する予定です。

また、移行中に IPv6 トラフィックが IPv4 ネットワークでブラックホール化してはなりません。アーキテクトはどの 2 つのソリューションを選択する必要がありますか? (2 つ選択してください。)

- A. C1 および C2 のアドレス ファミリ ipv6 でマルチトポロジが有効になっています
- B. すべてのルーターでメトリックスタイルの遷移が有効になっています
- C. E1 および E2 のアドレス ファミリ ipv6 でマルチトポロジが有効になっています
- D. C1 と C2 でメトリック スタイルの遷移が有効になっています
- E. E1 と E2 でメトリック スタイルの遷移が有効になっています

正解: C,E (コメントを发表する)

The design must enable IPv6 multi-topology behavior at the appropriate edge routers and move metric handling toward transition support where the current IS-IS network still uses narrow metrics. Multi-topology IS-IS is important during IPv4-to-IPv6 migration because IPv4 and IPv6 reachability do not always exist on identical links at every stage of the migration. If IPv6 follows the IPv4 topology when some devices or paths are not IPv6-ready, traffic can be blackholed. Multi-topology allows independent IPv4 and IPv6 SPF calculations, which prevents IPv6 traffic from being forwarded through IPv4-only portions of the network.

The metric-style transition setting is used when migrating from narrow to wide metrics, allowing interoperability while routers are upgraded and the design scales. Because the requirement says IPv6 areas will grow and the current network uses narrow metrics, the selected edge devices must support the IPv6 topology and metric transition behavior shown in the exhibit. Reference topics: IS-IS IPv6, multi-topology IS- IS, narrow metrics, wide metric transition, IPv6 migration, blackhole prevention.

有効的な**300-420J**問題集はJPNTTest.com提供され、**300-420J**試験に合格することに役に立ちます！JPNTTest.comは今最新**300-420J**試験問題集を提供します。JPNTTest.com 300-420J試験問題集はもう更新されました。ここで**300-420J**問題集のテストエンジンを手に入れます。最新版のアクセス、<https://www.jpntest.com/shiken/300-420J-mondaishu> **457**問、**30%ディスカウント**、特別な割引コード：**JPNshiken**」

質問: 92

ネットワーク エンジニアは、IP アドレス 172.16.15.12/32 のループバック インターフェイスを構成するスクリプトを準備します。会社のセキュリティ ポリシーに準拠するには、'Content-type' を使用します。

スクリプトに application/yang-data+json」が追加されました。ネットワーク デバイスへの接続を保護する必要があります。

この要件を満たすために、ネットワーク エンジニアはどのコード スニペットを使用する必要がありますか？





- A. オプションA
- B. オプションB
- C. オプションC
- D. オプションD

正解: **A** ([コメントを发表する](#))

The correct code snippet must use a secure HTTPS-based RESTCONF transaction with the YANG JSON media type. The requirement explicitly states that the content type is application/yang-data+json and that the connection to the network device must be secured. In Cisco IOS XE programmability, RESTCONF uses HTTP methods against YANG-modeled resources, and secure deployments use HTTPS rather than clear-text HTTP. A loopback interface with address 172.16.15.12/32 would be configured by sending a JSON payload that follows the selected YANG model structure, commonly an IETF or Cisco native interface model, with the correct content-type and authentication headers. The key security indicator in the option must therefore be HTTPS/TLS, not an insecure HTTP URL or a non-YANG content type. NETCONF over SSH would also be secure, but the stated content-type points to RESTCONF with JSON encoding. Option A is the valid snippet because it aligns the payload format with a secure transport. Reference topics: RESTCONF, YANG JSON encoding, application/yang-data+json, HTTPS, IOS XE programmability.

質問: 93

複数の支店を持つグローバル組織は、オーバーレイVPNソリューションを設計するためにネットワークアーキテクトを雇いました。ブランチは頻繁に相互に通信します。顧客は、将来さらにブランチを追加することを期待しています。顧客のセキュリティ要件を満たすために、アーキテクトは動的IPsecトンネルを使用してトラフィック保護を提供することを計画しています。アーキテクトはどのソリューションを選択する必要がありますか？

- A. DMVPN
- B. EasyVPN
- C. GETVPN
- D. L2TP

正解: **A** ([コメントを发表する](#))

DMVPN is the correct overlay VPN solution for dynamic, secure, scalable branch-to-branch communication.

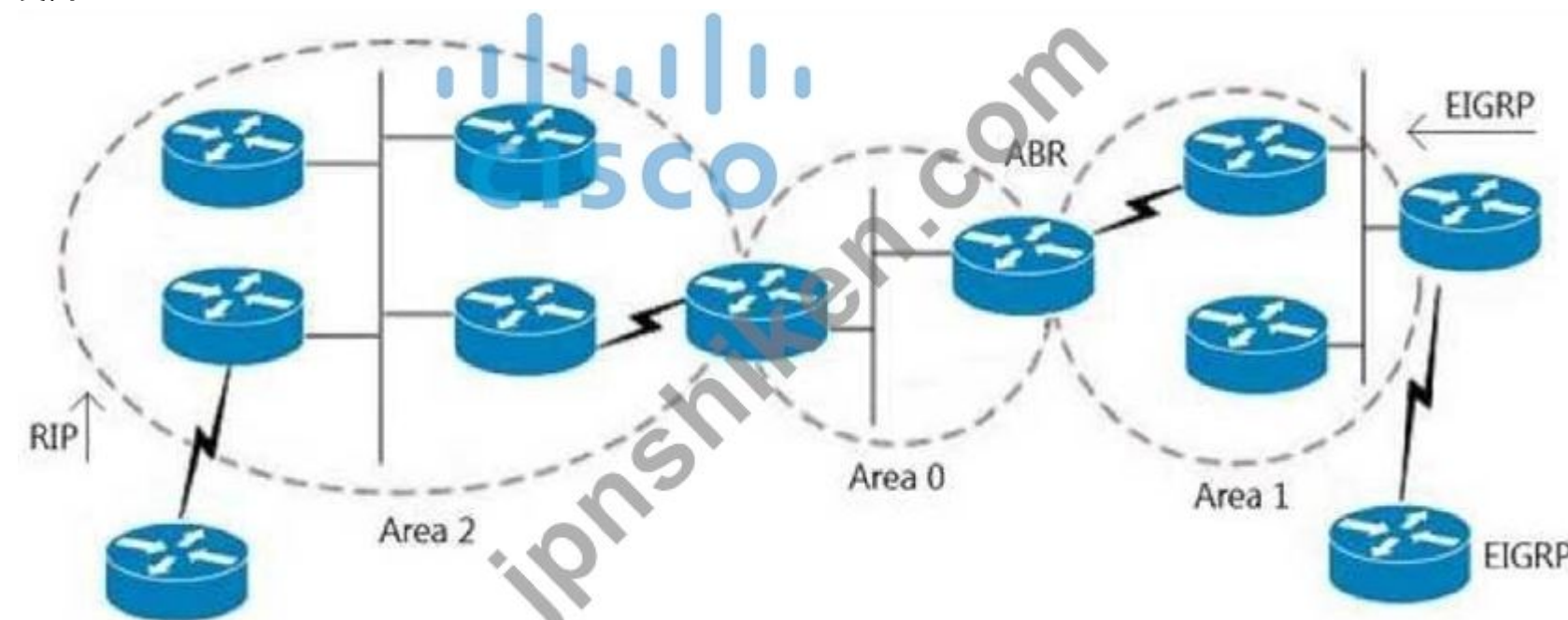
Cisco Dynamic Multipoint VPN uses multipoint GRE, NHRP, and IPsec to support dynamic tunnel creation without manually defining a static tunnel between every branch pair. That design is well suited to a global organization where branches communicate frequently and more branches will be added later. DMVPN allows spokes to register with hubs and, depending on the phase design, establish dynamic spoke-to-spoke tunnels for direct traffic while maintaining encrypted protection with IPsec. EasyVPN is more oriented to remote- access or simplified client/server VPN deployments and is not the scalable branch mesh solution described.

GETVPN is strong for private WAN encryption but does not build dynamic IPsec tunnels over an overlay in the same manner and assumes a trusted private routing core. L2TP is not the correct enterprise dynamic branch overlay. The architect should choose DMVPN and select the appropriate phase, routing protocol, summarization strategy, and IPsec profile based on branch scale and resource limits.

Reference topics:

DMVPN, NHRP, mGRE, IPsec protection, spoke-to-spoke dynamic tunnels.

質問: 94



図を参照してください。エンジニアがクライアント用の OSPF ネットワークを設計しています。要件では、エリア 1 のルータは、RIP ドメインで発生したルートを除き、EIGRP を含むネットワークに属するすべてのルートを受信する必要があります。エンジニアはどのようなアクションを取る必要がありますか。

- A. エリア 1 を NSSA にします。
- B. エリア 1 をスタブにします。
- C. エリア 1 を標準 OSPF エリアにします。
- D. エリア 1 ルータをエリア 0 の一部にします。

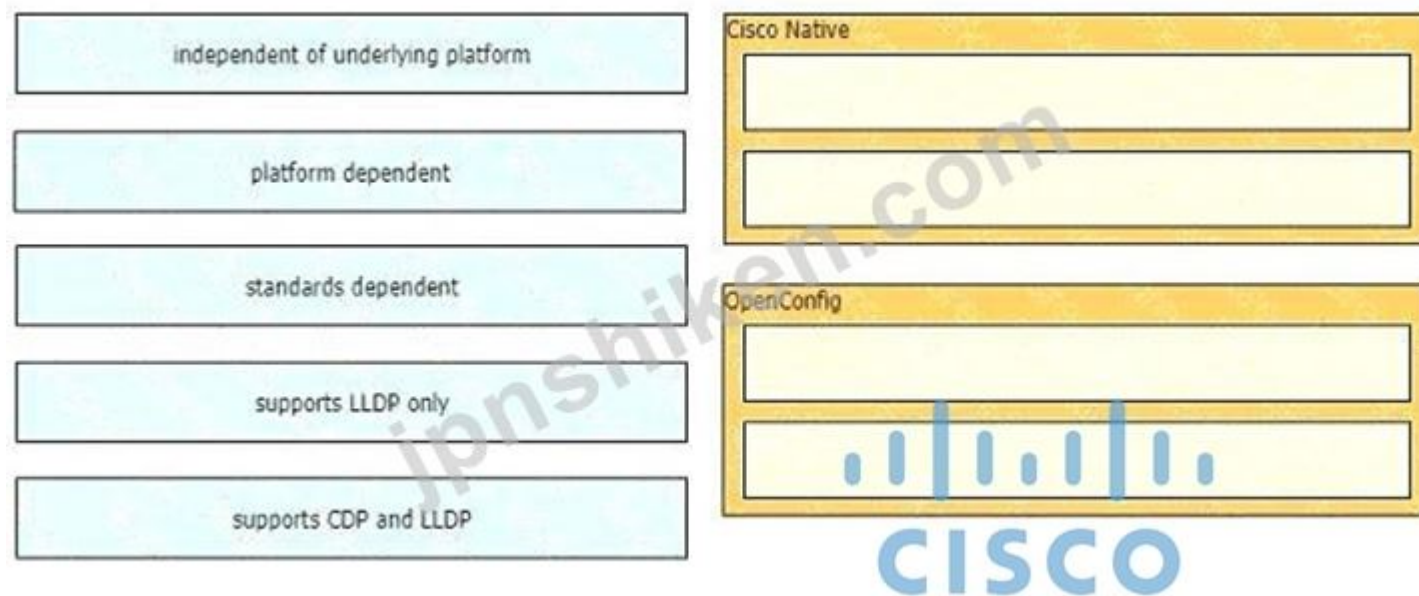
正解: [\(正解を表示します\)](#)

Area 1 should be configured as an NSSA when the design must allow selected external routes, such as EIGRP- originated routes, while controlling other external routes such as those from the RIP domain. A normal OSPF area receives Type 5 external LSAs, so it would not inherently block the unwanted RIP-originated external routes. A stub area blocks Type 5 external LSAs, but it also cannot contain an ASBR that injects external routes in the normal way. An NSSA is designed for this exact middle ground: it restricts external LSAs from the backbone while allowing external routes to be injected inside the area as Type 7 LSAs, which are translated by the ABR when needed. This allows Area 1 to receive the necessary internal and selected redistributed information while excluding external routes originated elsewhere. Making the routers part of area 0 would not solve the external-route filtering requirement and would weaken the intended area structure.

Therefore, NSSA provides the required balance of route visibility and external-route control. Reference topics: OSPF NSSA, Type 5 LSAs, Type 7 LSAs, external route control, area design.

質問: 95

左側の特性を、右側で説明されている YANG モジュールにドラッグ アンド ドロップします。すべてのオプションが使用されるわけではありません。



正解:



Explanation:

The YANG module mapping follows the distinction between standards-based, vendor-neutral, and vendor- native data models. IETF YANG models are standards-driven and are appropriate when the design goal is protocol-level interoperability across vendors for common network functions. OpenConfig models are also vendor-neutral, but they are developed by the OpenConfig community and commonly use a structured separation between intended configuration and operational state. Cisco native models expose platform- specific Cisco IOS XE, IOS XR, or NX-OS capabilities and are often the most complete option when a feature is unique to a Cisco operating system. The selected drag-and-drop answer aligns each characteristic with the model family that owns that behavior. In a professional automation design, the engineer should not choose a model only because it is familiar. The correct choice depends on feature coverage, device support, platform release, schema stability, and whether the workflow must operate consistently across vendors. For multivendor intent, choose IETF or OpenConfig first; for Cisco-specific feature depth, use Cisco native models. Reference topics: YANG, IETF models, OpenConfig, Cisco native models, NETCONF and RESTCONF automation.

質問: 96

エンジニアは、2つのホストが通信できるようにするネットワーク ソリューションを設計しています。1つのホストは会社 A のネットワーク内にあり、もう 1つのホストは会社 B のネットワーク内にあります。両社には、将来的に接続を追加する予定はありません。両社とも、単一の安全で暗号化されたインターネット接続を使用したいと考えていますが、構成はできるだけシンプルにする必要があります。エンジニアはどのネットワーク ソリューションを選択する必要がありますか？

A. EIGRPルーティングによる単一のDMVPN

- B. OSPFルーティングによるルーティングされたIPsecトンネル
- C. 静的ルーティングによるポリシーベースの IPsec トンネル
- D. BGPルーティングによるMPLS VPN提供サービス

正解: **C** ([コメントを發表する](#))

A policy-based IPsec tunnel with static routing is the simplest secure design for connecting two hosts in two companies when no additional future connectivity is planned. The requirement is narrow: one host in Company A must communicate with one host in Company B over a single encrypted Internet connection.

Policy-based IPsec is straightforward for this type of point-to-point encryption because crypto ACLs define exactly which interesting traffic should be protected. Static routing is sufficient because there is no complex topology, no dynamic path selection requirement, and no anticipated expansion. A DMVPN design would add unnecessary NHRP, multipoint GRE, and dynamic routing complexity. A routed IPsec tunnel with OSPF is useful when many networks or dynamic route changes are expected, but it is more complex than needed.

MPLS VPN with BGP would require provider service and does not meet the simple encrypted Internet connection requirement. Reference topics: site-to-site IPsec, policy-based VPN, static routing, point-to-point secure connectivity, WAN simplicity.

質問: **97**

SD-WANオーバーレイにおいて、オーバーレイ管理プロトコルはどのルートをアドバタイズしますか？

- A. プライマリ、バックアップ、ロードバランシング
- B. インターネット、MPLS、およびバックアップ
- C. 接頭辞、TLOC、およびサービス
- D. アンダーレイ、MPLS、オーバーレイ

正解: ([正解を表示します](#))

質問: **98**

ネットワークソリューションは、複数のインターネットサービスプロバイダーに接続する企業向けに設計されています。
どっち

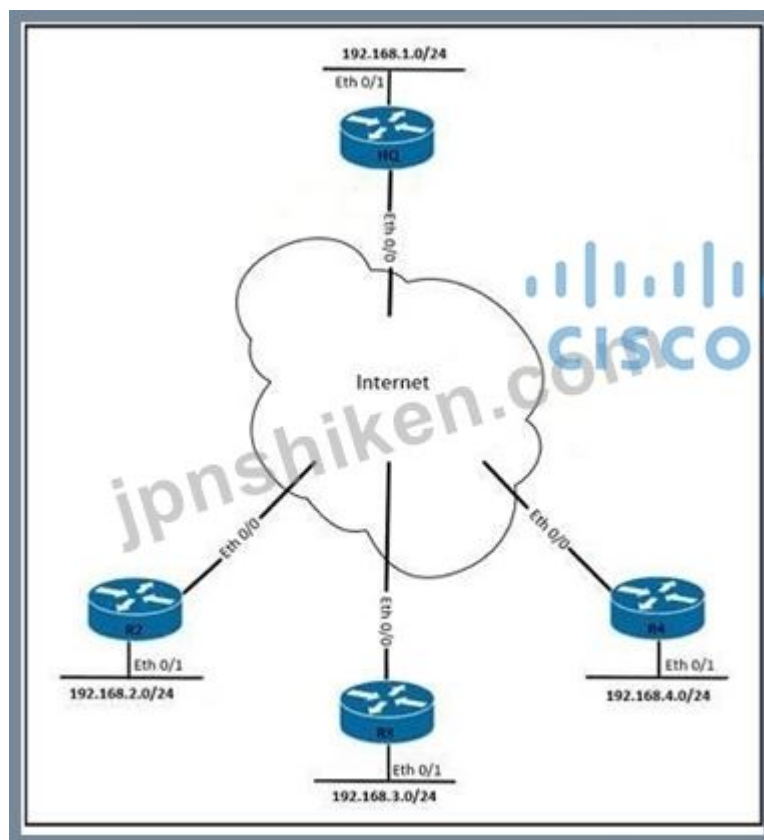
シスコ独自のBGPパス属性は、発信トラフィックフローに影響しますか？

- A. ローカル設定
- B. MED
- C. 重量
- D. ASパス
- E. コミュニティ

正解: ([正解を表示します](#))

Cisco weight is the Cisco-proprietary BGP path attribute used to influence outbound traffic selection on a local router. Weight is locally significant and is not advertised to BGP neighbors. A higher weight is preferred, and because weight is evaluated early in the Cisco BGP best-path algorithm, it is a direct way to make one exit path preferred over another on the router where the attribute is configured. Local preference also influences outbound traffic, but it is a well-known discretionary BGP attribute used inside an AS and is not Cisco proprietary. MED is primarily used to influence inbound path selection by suggesting to a neighboring AS which entry point should be preferred. AS-path prepending is also typically used for inbound traffic engineering by making a route appear less attractive to external peers. Communities can carry policy information, but they are not the Cisco-proprietary attribute in the list. Therefore, when the requirement specifically asks for the Cisco proprietary BGP path attribute that influences outbound traffic flow, the answer is weight.

質問: **99**



図を参照してください。顧客は、リモート ブランチと本社間の VoIP、ビデオ、FTP トラフィックの通信を保護するために、動的なサイト間 VPN ソリューションを導入したいと考えています。また、顧客はブランチが直接通信できるようにして、本社のトラフィックを削減したいと考えています。ソリューションでは、ブランチ ルーターの使用可能なメモリが限られていることを考慮する必要があります。これらの要件を満たす VPN ソリューションはどれですか。

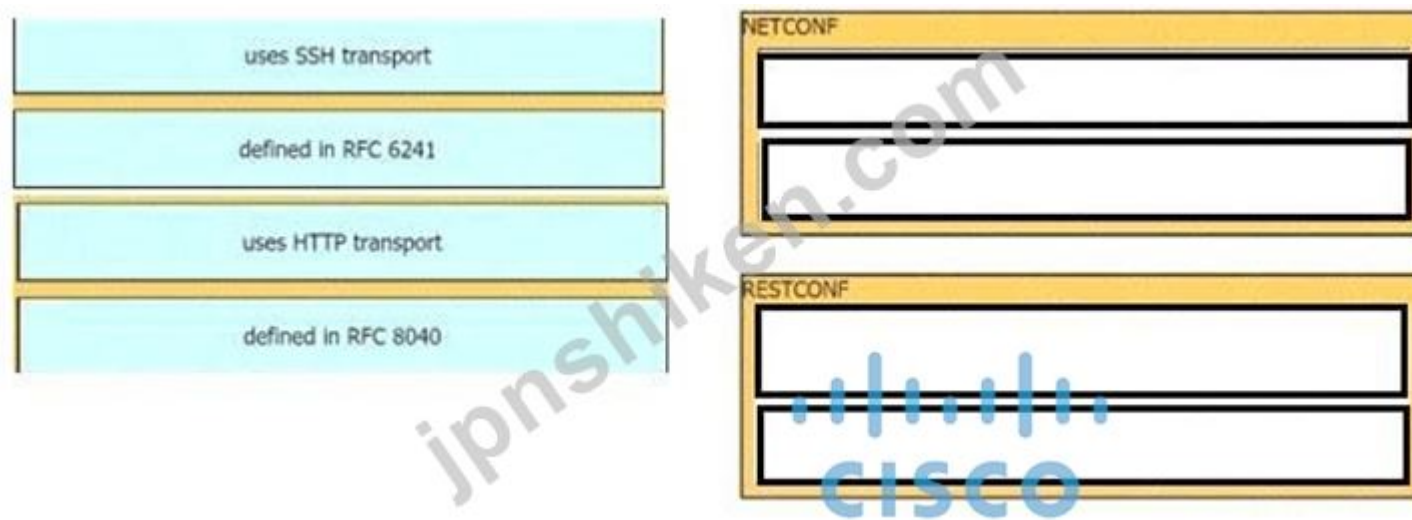
- A. DMVPN フェーズ 2 ハブ アンド スポーク設計
- B. DMVPN フェーズ 3 ハブ アンド スポーク設計
- C. DMVPN フェーズ 1 ハブ アンド スポーク設計
- D. DMVPN フェーズ 3 階層設計

正解: **B** ([コメントを发表する](#))

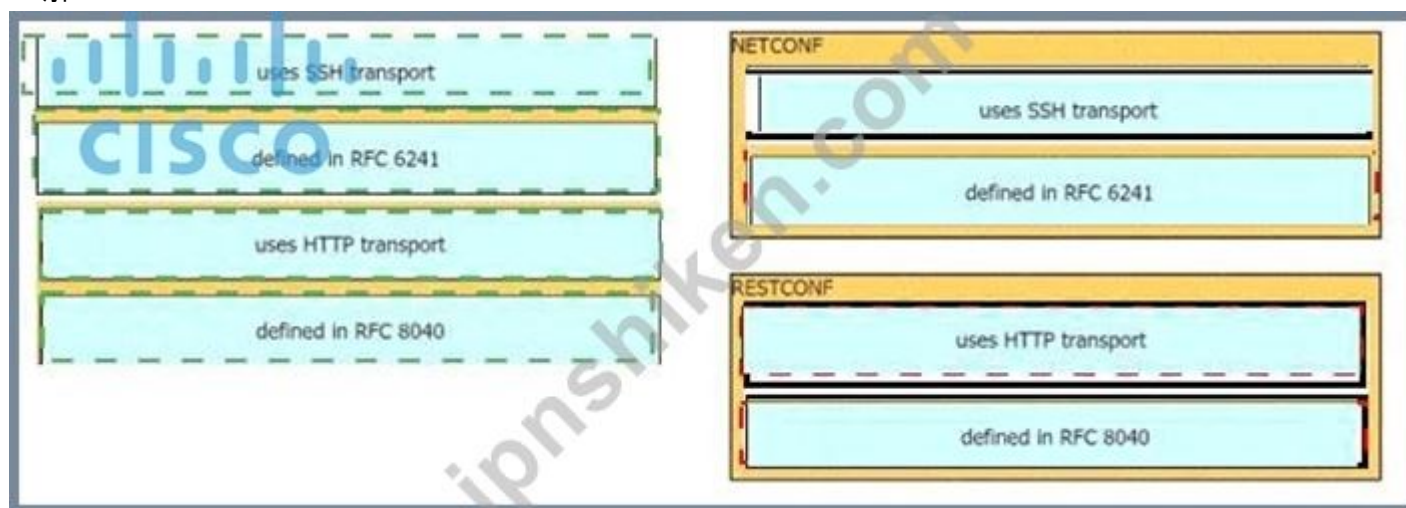
DMVPN Phase 3 hub-and-spoke is the best design for this requirement. The customer wants secure dynamic site-to-site VPN connectivity for voice, video, and FTP, and also wants branches to communicate directly instead of hairpinning all traffic through headquarters. DMVPN supports dynamic spoke-to-spoke tunnels using multipoint GRE, NHRP, and IPsec protection. Phase 3 is more scalable than Phase 2 because the hub can send NHRP redirect messages and spokes can install shortcut routes, reducing the need for every spoke to maintain detailed next-hop information for every other spoke. That is important because the branch routers have limited memory. Phase 1 is hub-and-spoke only and does not provide dynamic spoke-to-spoke forwarding. A hierarchical Phase 3 design can be useful at very large scale, but the question describes a headquarters and remote branch design rather than multiple tiers of hubs. Phase 3 hub-and-spoke satisfies direct branch communication, security, and scale. Reference topics: DMVPN Phase 3, NHRP redirect, spoke- to-spoke tunnels, IPsec encryption, branch WAN scalability.

質問: **100**

左側の特性を、右側で説明されている構成プロトコルにドラッグ アンド ドロップします。



正解:



Explanation:

The configuration-protocol mapping distinguishes NETCONF and RESTCONF behavior in model-driven programmability. NETCONF is an XML-based network configuration protocol that runs over a secure transport such as SSH and supports structured operations such as get-config, edit-config, lock, unlock, commit, and datastore manipulation. It is well suited to transactional configuration workflows where candidate and running datastores are supported. RESTCONF exposes YANG-modeled data through a REST-style interface over HTTP or HTTPS, using familiar methods such as GET, POST, PUT, PATCH, and DELETE. It can encode data in XML or JSON depending on headers and implementation. Both protocols use YANG models to define the structure and semantics of configuration and operational data, but they differ in transport style, operation model, and tooling. The selected mapping therefore places NETCONF characteristics with the protocol that provides RPC-style operations and datastore handling, while RESTCONF receives HTTP method and URI-based characteristics. In design, the choice depends on controller support, device operating system, required transaction semantics, security policy, and developer tooling. Reference topics: NETCONF, RESTCONF, YANG, datastores, XML and JSON encoding.

質問: 101

Cisco SD-Accessオーバーレイ設計の主な目的は何ですか？

- A. サポートチームによるネットワーク管理とトラブルシューティングを簡素化するため
- B. ユーザーサービスの高い可用性と耐障害性を確保するため
- C. SD-Accessオーバーレイサービスとのシームレスな統合を可能にするため
- D. インフラストラクチャのネットワーク可視性と監視を強化する

正解: ([正解を表示します](#))

The main purpose of the Cisco SD-Access overlay design is to enable the virtualized fabric services that run over the physical underlay. Cisco describes the overlay as the logical network created on top of the routed underlay, using virtual networks to provide segmentation and policy separation. SD-Access uses LISP for endpoint-to-location mapping and VXLAN for data-plane encapsulation, allowing Layer 2

and Layer 3 services to be delivered across a stable Layer 3 transport. Option C is the best available answer because the overlay exists to integrate and carry SD-Access services such as virtual networks, endpoint mobility, macrosegmentation, and security group policy. Simplified troubleshooting and visibility are management and assurance outcomes, not the overlay's main technical purpose. High availability depends on the underlay topology, node redundancy, and control-plane design, but it is not the primary definition of overlay design. Reference topics: SD-Access overlay, virtual networks, LISP, VXLAN, segmentation, fabric services.

質問: 102

Cisco SD-Accessアーキテクチャにおけるコントロールプレーンノードの役割は何ですか？

- A.** 有線エンドポイントをSD-Accessファブリックに接続するファブリックデバイス
- B.** エンドポイントとデバイスの関係を管理するマップシステム
- C.** APとワイヤレスエンドポイントをSD-Accessファブリックに接続するファブリックデバイス
- D.** 外部レイヤー3ネットワークを管理するマップシステム

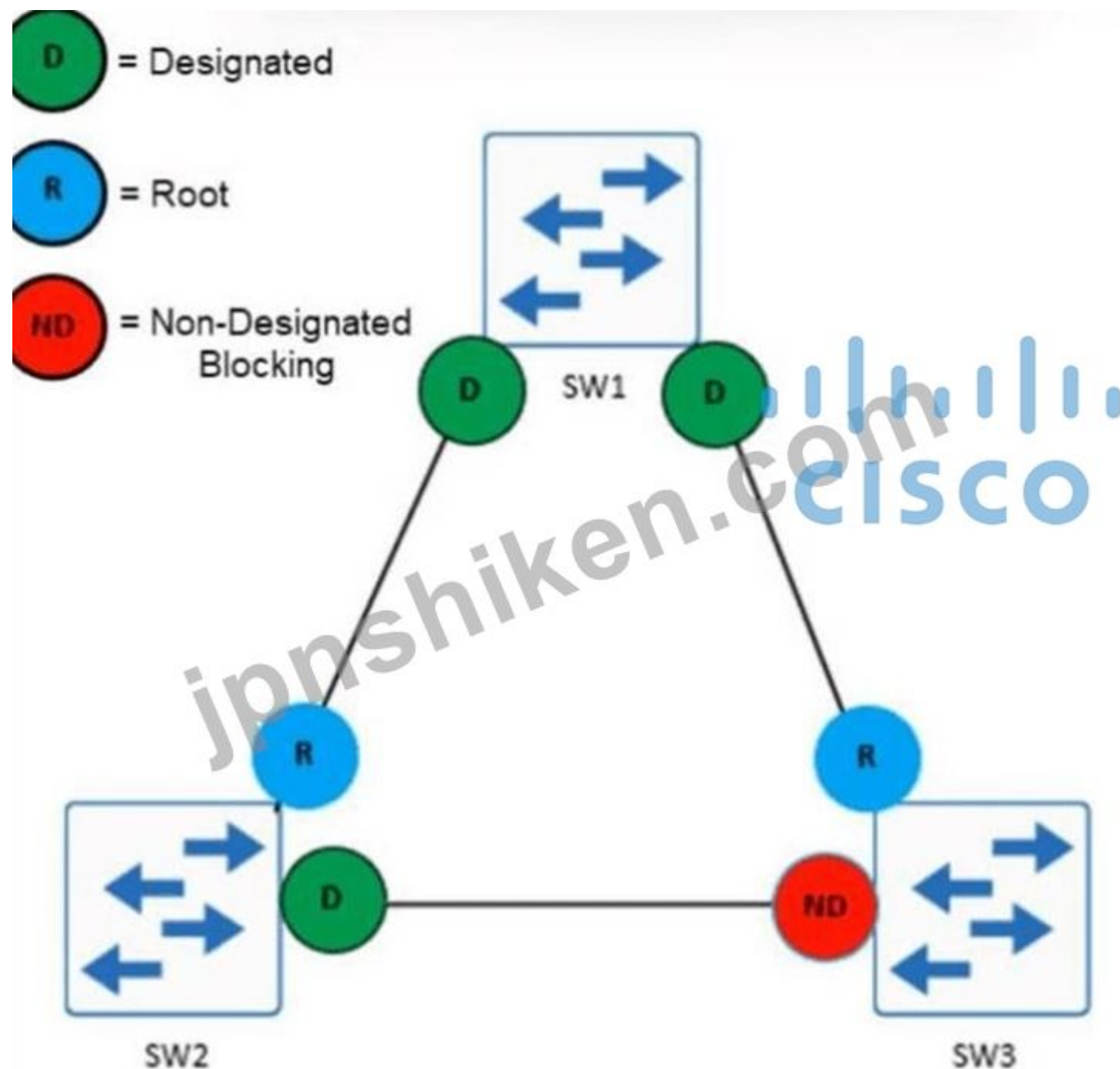
正解: ([正解を表示します](#))

The control-plane node in Cisco SD-Access acts as the map system that manages endpoint-to-device relationships. Cisco SD-Access documentation describes control-plane nodes as maintaining the host database that maps endpoint identifiers to the fabric edge devices where those endpoints are attached. This mapping function is based on LISP control-plane behavior. When an endpoint appears on an edge node, the edge registers the endpoint information with the control plane. When another fabric device needs to send traffic to that endpoint, it queries the control-plane node to resolve the endpoint location. A fabric edge node is the device that connects wired endpoints to the fabric, so option A describes an edge role. Wireless endpoints may attach through fabric AP and fabric WLC integration, but that is not the control-plane node's role.

External Layer 3 networks are connected through border nodes, not through the control-plane node.

Therefore, the correct role is endpoint-to-device mapping. Accurate control-plane design is essential for fabric scale, endpoint mobility, and fast location resolution.

質問: 103



図を参照してください。アーキテクトが顧客向けにレイヤー 2 ネットワークを設計しています。ネットワークではスパニングツリー プロトコルを使用します。SW1 と SW2 間のリンク障害時には、可能な限り最短の収束時間が求められます。アーキテクトはどのソリューションを選択する必要がありますか？

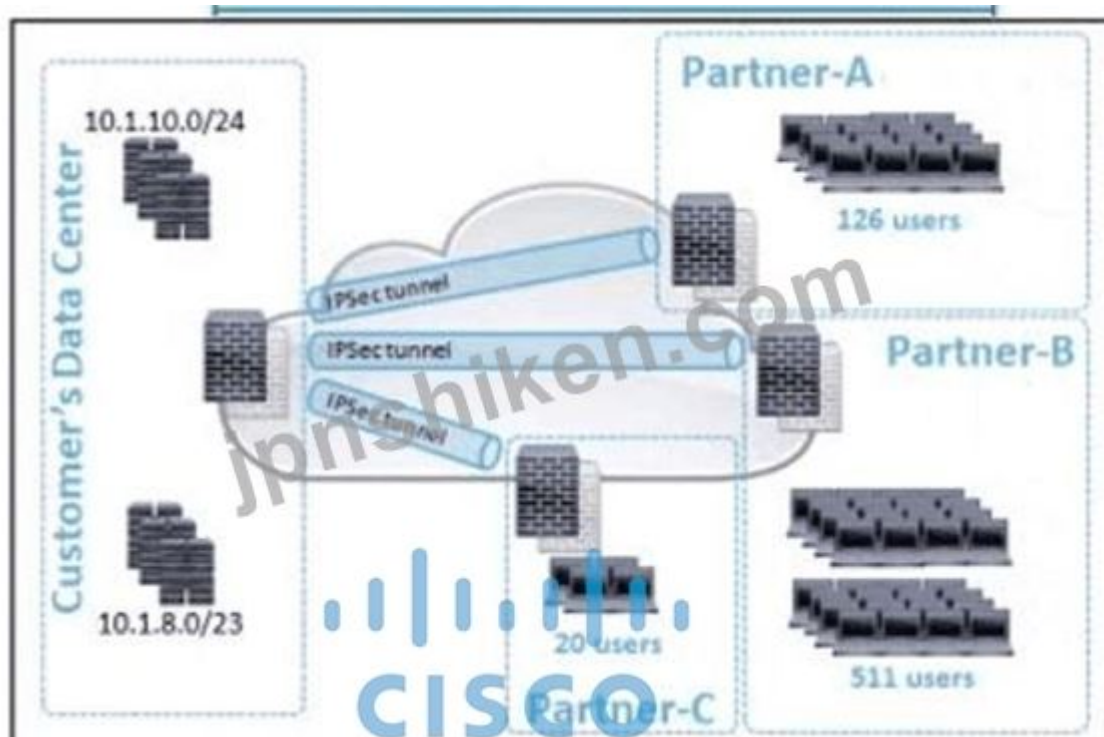
- A. ループガード
- B. アップリンクファースト
- C. ポートファスト
- D. バックボーンファースト

正解: [\(正解を表示します\)](#)

UplinkFast is the correct spanning-tree feature for faster convergence after an uplink failure in a classic STP access-layer design. UplinkFast is designed for access switches that have redundant uplinks toward the distribution layer. When the primary root port fails, UplinkFast allows an alternate blocked uplink to transition rapidly to forwarding, reducing the outage compared with standard 802.1D timer-based convergence. PortFast is used on edge ports connected to end hosts, not on switch-to-switch uplinks. Loop Guard protects against unidirectional failures or missing BPDUs on non-designated ports, but it does not accelerate recovery from an access uplink failure. BackboneFast helps when an indirect link failure is detected elsewhere in the topology, but the described failure is directly between SW1

and SW2. The design objective is the fastest possible convergence during that direct uplink loss, which matches UplinkFast in a traditional STP environment. Reference topics: Spanning Tree Protocol, UplinkFast, access-layer redundancy, root port failover, Layer 2 convergence.

質問: 104



図を参照してください。顧客は、データセンターに3つの新しいVPNパートナー接続をオンボードすることを計画しています。

新しいサブネットは既存のデータセンターネットワークと重複してはならず、サブネットサイズは必要以上に大きくしてはなりません。お客様はこの設計に 10.1.8.0/21 を専用に使いました。Ho1// これらの要件を満たすためにサブネットを分割する必要がありますか？

- Partner-A: 10.1.11.0/25
Partner-B: 10.1.12.0/22
Partner-C: 10.1.11.128/27
- A. Partner-A: 10.1.11.0/24
Partner-B: 10.1.13.0/23
Partner-C: 10.1.12.128/26
- B. Partner-A: 10.1.11.0/25
Partner-B: 10.1.13.0/23
Partner-C: 10.1.11.128/27
- C. Partner-A: 10.1.12.0/25
Partner-B: 10.1.13.0/22
Partner-C: 10.1.12.128/26
- D. Partner-A: 10.1.12.0/25
Partner-B: 10.1.13.0/22
Partner-C: 10.1.12.128/26

正解: ([正解を表示します](#))

Option B is retained because the subnetting task depends on the exhibit options and the chosen plan must divide 10.1.8.0/21 into nonoverlapping partner VPN subnets without allocating more address space than required. The design principle is variable-length subnet masking: allocate the smallest subnet that satisfies each partner requirement, then place the subnets on valid prefix boundaries inside the reserved aggregate.

Cisco addressing design emphasizes conserving address space, avoiding overlap, and preserving summarization where possible. A /21 provides 2048 total addresses, so the designer can carve several smaller partner ranges while keeping the entire allocation summarizable as 10.1.8.0/21 toward the data-center core.

Incorrect options would either waste addresses, overlap existing ranges, place subnets on invalid boundaries, or allocate blocks larger than necessary. Option B is the mapping that satisfies the exhibit constraints while keeping the design clean for routing, firewall policy, and VPN onboarding. Reference topics: IPv4 VLSM, subnet boundary alignment, partner VPN addressing, summarization, nonoverlapping address design.

質問: 105

アーキテクトは、企業ネットワーク機器を管理するための計画を設計する必要があります。その設計には、以下の点を考慮する必要があります。

- * すべてのネットワーク機器に専用の管理インターフェースがあるわけではありません
- * すべてのデバイスのすべてのIP対応インターフェースに到達可能である必要があります
- * 暗号化は、サポートしているすべてのデバイスで使用する必要があります

建築家はどちらの解決策を選択すべきか？

- A. KVMサーバー
- B. インバンド
- C. 帯域外
- D. ターミナルサーバー

正解: ([正解を表示します](#))

In-band management is the correct design because the management plan must reach all IP-enabled interfaces, and not every device has a dedicated management port. In-band management uses the production IP network for administrative access, allowing management stations to reach loopbacks, SVIs, routed interfaces, or front-panel interfaces through normal routing. Cisco management guidance distinguishes in-band management from out-of-band designs that depend on a separate physical management network or console path. Because the question requires reachability to all IP-enabled interfaces, in-band is the only listed model that naturally supports that scope. Security must still be enforced with encrypted protocols where supported, such as SSH and HTTPS, plus ACLs, AAA, and management-plane protection where appropriate. A terminal server provides console access, not routed IP management to every interface. A KVM server is a server-management tool, not a network-device management architecture. Out-of-band management is preferred for isolation, but it does not fit devices without dedicated management access. Reference topics: in-band management, secure management protocols, SSH, HTTPS, network management design.

質問: 106

```
Switch(config)# interface gig0/2
Switch(config-if)# switchport port-security
Switch(config-if)# switchport port-security maximum 1
Switch(config-if)# switchport port-security mac-address 00-d0-ba-11-21-31
Switch(config-if)# switchport port-security violation shutdown
Switch(config-if)#end
```

展示を参照してください。Cisco Catalyst スイッチは、インターフェイス gig0/2 で 1 つの MAC アドレスのみを手動で学習するように設定されています。スイッチ ポートに接続されているデバイスを動的に学習するには、どのコマンドを実行する必要がありますか。

- ☐ `switchport port-security mac-address auto`
- ☐ `switchport port-security mac-address sticky`
- ☐ `switchport port-security aging`
- ☐ `switchport port-security maximum 1 auto`

CISCO

- A. オプションA
- B. オプションB
- C. オプションC
- D. オプションD

正解: ([正解を表示します](#))

The selected command is correct because the port is already restricted to one manually learned MAC address, and the requirement is to let the switch dynamically learn the connected endpoint while retaining port-security control. In Cisco Catalyst port security, a static secure MAC address is manually configured, while a sticky secure MAC address is dynamically learned and then added to the running configuration. This design is useful on access ports where the administrator wants the operational convenience of dynamic learning but still wants the switch to enforce the maximum secure-address limit. Once the endpoint MAC is learned, subsequent frames from different source MAC addresses can trigger the configured violation action, such as protect, restrict, or shutdown. The command must therefore enable sticky learning on the secured interface rather than simply changing the maximum, disabling port security, or configuring a normal switching feature. This is an access-layer security function, not a routing or QoS function. Reference topics: Catalyst port security, secure MAC addresses, sticky learning, access port hardening, Layer 2 campus design.

有効的な**300-420J**問題集はJPNTTest.com提供され、**300-420J**試験に合格することに役に立ちます！JPNTTest.comは今最新**300-420J**試験問題集を提供します。JPNTTest.com 300-420J試験問題集はもう更新されました。ここで**300-420J**問題集のテストエンジンを手に入れます。最新版のアクセス、<https://www.jpntest.com/shiken/300-420J-mondaishu> **457**問、**30%ディスカウント**、特別な割引コード：**JPNshiken**」

質問: 107

ある企業は、既存のサーバー上でネットワーク内におけるIPv6アプリケーションのパイロットプロジェクトを実施する必要があり、移行戦略を検討しています。単一のVLAN内に収まるこのパイロットプロジェクトは、レイヤ2スイッチとレイヤ3スイッチで構成される2つのサイトからなるデータセンター環境にまたがる必要があります。このパイロットプロジェクトにおける主要な考慮事項は何でしょうか？

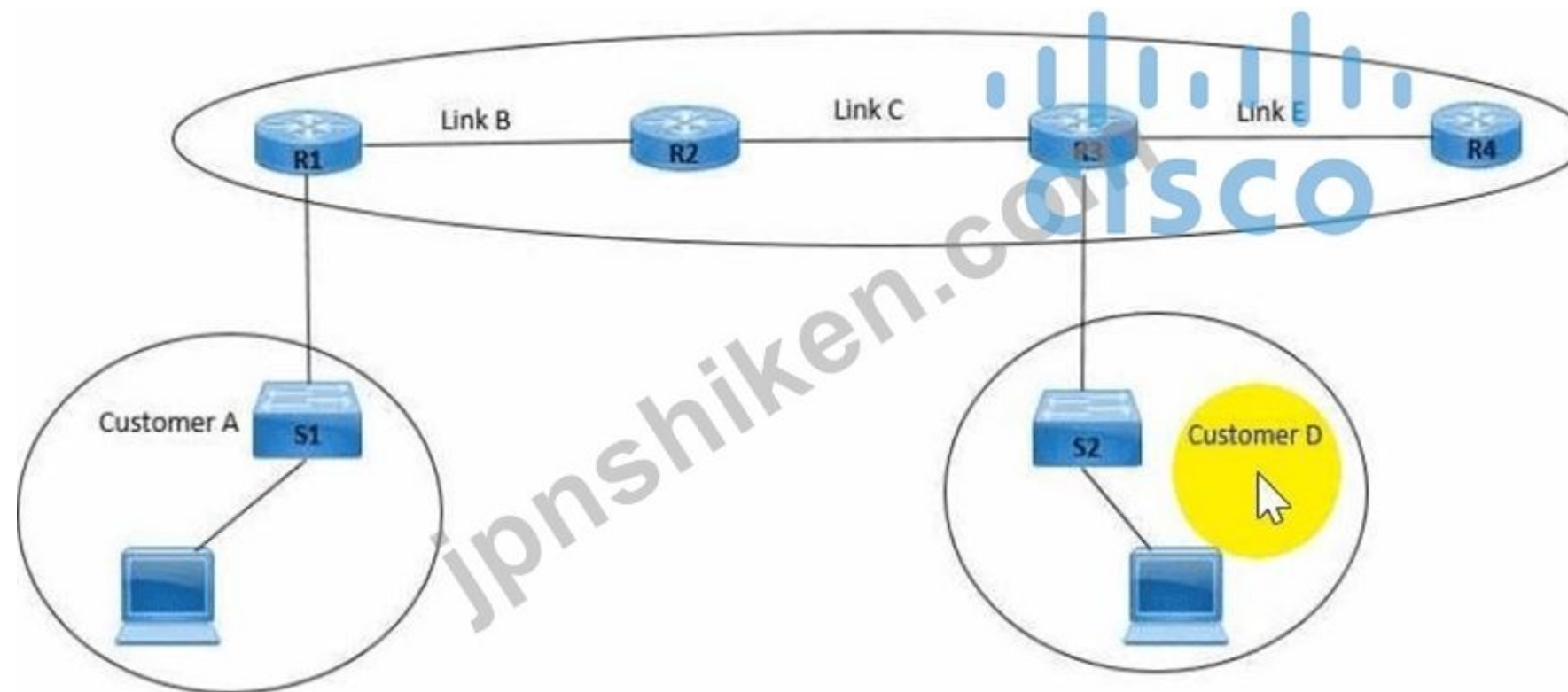
- A. データセンターネットワークをプロビジョニングする各データセンター内のレイヤ2およびレイヤ3スイッチは、デュアルスタッキングをサポートする必要があります。
- B. パイロットに参加する各データセンター内のホストは、デュアルスタッキングをサポートする必要があります。
- C. VLAN をプロビジョニングする各データセンター内のレイヤ2スイッチは、デュアルスタッキングをサポートする必要があります。
- D. ネットワークをプロビジョニングする各データセンター内のレイヤ3スイッチは、デュアルスタッキングをサポートする必要があります。

正解: ([正解を表示します](#))

The primary consideration is that the participating hosts must support dual stack. The pilot is contained within a single VLAN, so the application traffic remains inside the same Layer 2 broadcast domain across the dual-site data center environment. Layer 2 switches forward Ethernet frames and do not need to understand the IPv6 header to carry IPv6 traffic inside the VLAN. Layer 3 switches need IPv6 capability only where IPv6 routing, gateways, ACLs, or services are required. For this pilot, the essential requirement is that the servers and endpoints running the IPv6 application can operate with both IPv4 and IPv6 enabled while the existing environment continues to support IPv4. Dual stack allows IPv6 testing without removing IPv4 reachability or forcing translation. Requiring every Layer 2 and Layer 3

switch in both data centers to be dual-stack overstates the design requirement. The hosts participating in the pilot are the critical dependency. Reference topics: IPv6 migration, dual stack, Layer 2 VLAN transport, data-center pilot design, IPv4 and IPv6 coexistence.

質問: 108



展示を参照してください。アーキテクトは、172.20.0.0/16ネットワークを使用してIPv4プランを設計しています。設計では、サブネットの数を最大化し、無駄なIPアドレスの数を最小限に抑える必要があります。さらに、プランでは、これらの顧客とリンクにサブネットを割り当てる必要があります。

* 125台のホストをサポートする顧客A

* 62台のホストをサポートする顧客D

* リンク B C. および E

これらの要件を満たす 2 つの構成セットはどれですか? (2 つ選択してください)

Customer A - 172.20.0.128/25
Customer D - 172.20.1.0/26

A.

Link B - 172.20.1.70/30
Link C - 172.20.1.74/30
Link E - 172.20.1.78/30

B.

Customer A - 172.20.1.0/24
Customer D - 172.20.2.64/26

C.

Link B - 172.20.1.68/30
Link C - 172.20.1.72/30
Link E - 172.20.1.76/30

D.

Link B - 172.20.2.132/30
Link C - 172.20.2.136/30
E. Link E - 172.20.2.140/30

正解: ([正解を表示します](#))

The selected configuration sets are correct because the addressing plan must use VLSM to maximize subnets and minimize wasted addresses. Customer A requires 125 hosts, so it needs a /25 subnet, which provides 126 usable host addresses. Customer D requires 62 hosts, so it needs a /26 subnet, which provides 62 usable addresses exactly. Point-to-point or small WAN links B, C, and E should use the smallest appropriate link subnets shown in the exhibit, typically /30 for IPv4 point-to-point links when two usable addresses are required. A correct plan must allocate the largest blocks first, align every subnet on the proper boundary, avoid overlap, and avoid assigning a larger-than-needed subnet to any customer or link. Options A and C are the two valid combinations that satisfy those rules in the exhibit. Any option that places a /25 or /26 on the wrong boundary, overlaps with another block, or wastes larger subnets on links would fail the design objective. Reference topics: IPv4 VLSM, subnet boundary alignment, host capacity planning, point-to-point addressing, route summarization readiness.

質問: 109

企業は、最近、唯一の MPLS プロバイダーで障害が発生したため、WAN の可用性を強化する必要があります。提案されたソリューションは、迅速に導入でき、手頃な価格で、信頼性が高く、企業のプライマリ MPLS 接続のバックアップとして機能する必要があります。これらの要件を満たすソリューションはどれですか。

- A. インターネット接続を契約し、DMVPN を展開します。
- B. BFDエコーモードを展開し、プロバイダーPEをプローブします。
- C. 追加のWANルーターを導入し、フローティングスタティックルートを使用する
- D. 別の MPLS プロバイダーと契約し、GET VPN を展開します。

正解: **A** ([コメントを發表する](#))

Contracting an internet connection and deploying DMVPN is the most practical solution when the enterprise needs a fast, affordable, and reliable backup for a single MPLS provider. DMVPN can use broadband internet transport to build encrypted GRE/IPsec tunnels between sites, providing a backup WAN path without waiting for another private MPLS provider circuit. It supports dynamic routing and can be deployed incrementally, which makes it suitable after a recent outage when time and cost matter. BFD echo mode toward a provider PE can improve failure detection, but it does not create an alternate transport if the only MPLS provider has an outage. Adding a WAN router and a floating static route also fails unless another path exists. Contracting a second MPLS provider and deploying GETVPN may be robust, but it is slower and more expensive than the requested quick and affordable backup. Reference topics: DMVPN backup WAN, MPLS resiliency, internet VPN, IPsec, branch availability, rapid WAN deployment.

質問: 110

IPsec VPNが設計されている場合、IPマルチキャストのサポートが必要な場合の固有の要件は何ですか？

- A. GREまたはVTIによるトラフィックのカプセル化
- B. トランスポートモードを使用したIPsec転送
- C. ヘッドエンドの追加帯域幅
- D. トンネルモードを使用したIPsec転送

正解: ([正解を表示します](#))

Native policy-based IPsec protects unicast IP traffic between defined proxy identities, but it does not inherently provide a multiprotocol tunnel interface that supports multicast routing protocol packets. Many enterprise VPN designs require multicast for applications, routing protocols, or service discovery. To support multicast across an encrypted VPN, the design commonly uses GRE over IPsec or a tunnel interface such as VTI where the tunnel provides a logical point-to-point or multipoint interface and IPsec secures the encapsulated traffic. GRE can encapsulate multicast and broadcast traffic, allowing dynamic routing protocols and multicast packets to traverse the VPN. IPsec then encrypts the GRE or tunnel traffic for confidentiality and integrity. Transport mode by itself does not solve multicast forwarding, and tunnel mode still requires suitable encapsulation or tunnel-interface behavior for multicast support. Additional headend bandwidth may be required in large VPN designs, but it is not the unique technical requirement for multicast. Therefore, when IPsec VPNs must support IP multicast, the design must include GRE or VTI-style encapsulation.

質問: 111

ある会社は、2つの新しい支店を開設し、その地域に2a01:c30:16:7009::3800/118IPv6ネットワークを割り当てることを計画しています。各ブランチには、最大200のホストを収容できる容量が必要です。会社はどちらの2つのネットワークを使用する必要がありますか？ 2つ選択してください。)

- A. 2a01:0c30:0016:7009::3a00/120
- B. 2a01:0c30:0016:7009::3b00/121
- C. 2a01:0c30:0016:7009::3a80/121
- D. 2a01:0c30:0016:7009::3b00/120
- E. 2a01:0c30:0016:7009::3c00/120

正解: ([正解を表示します](#))

The branch subnets must each support up to 200 hosts, so /120 is the correct IPv6 prefix length from the available options. A /120 provides 256 IPv6 addresses, which satisfies the 200-host requirement.

A /121 provides only 128 addresses, which is not enough. The parent block is 2a01:0c30:0016:7009::3800/118. A

/118 has 10 host bits, so it contains 1024 addresses and spans from ::3800 through ::3bff. The valid /120 boundaries inside that parent are ::3800/120, ::3900/120, ::3a00/120, and ::3b00/120. From the listed choices,

2a01:0c30:0016:7009::3a00/120 and 2a01:0c30:0016:7009::3b00/120 both fall inside the parent allocation and provide enough address capacity. The /121 options are too small, and ::3c00/120 is outside the /118 parent range. This is a straightforward IPv6 prefix-boundary and capacity calculation. Reference topics: IPv6 subnetting, prefix length, address block boundaries, branch address planning.

質問: 112

フラグメンテーションの問題を回避するために、Cisco SD-WANはどの方法を使用しますか？

- A. PMTUDが使用されます。
- B. トラフィックはDFビットセットでマークされます。
- C. ジャンボフレームが有効になります。
- D. アクセス回線は1600バイトのMTU設定で構成されています。

正解: ([正解を表示します](#))

Cisco SD-WAN uses Path MTU Discovery to avoid fragmentation problems across the overlay. SD-WAN traffic is encapsulated and encrypted, which adds overhead beyond the original packet size. If the effective path MTU is smaller than the encapsulated packet, fragmentation or packet loss can occur, especially when the DF bit is set or when intermediate devices block fragmentation-related ICMP messages. PMTUD allows devices to discover the largest packet size that can traverse the path without fragmentation and adjust forwarding behavior accordingly. Marking traffic with DF alone does not solve the issue; it may expose the problem by causing packets to be dropped when too large. Jumbo frames are not a universal WAN solution because service-provider, Internet, and broadband paths often do not support them consistently. Configuring every access circuit with a fixed 1600-byte MTU is not a reliable design assumption. The correct design is to rely on PMTUD and validate tunnel overhead, device MTU, TCP MSS adjustment where needed, and ICMP handling across transports. Reference topics: Cisco SD-WAN MTU, PMTUD, tunnel overhead, fragmentation avoidance, IPsec encapsulation.

質問: 113

ある顧客は、以下の高レベル設計要件に基づいて、企業内でWoL(Wake-on-LAN)を導入する計画を立てています。

DHCPサービスが利用可能である必要があります。

クライアントのBIOS設定でWoLを有効にする必要があります。

クライアントはオンラインになるとIPアドレスを取得します。

レイヤ2スイッチでは、スパニングツリーPortFastが有効になっています。

顧客が導入を成功させるためには、どの2つのソリューションを選択する必要がありますか？(2つ選択してください。)

- A. クライアントサブネットが存在するすべてのSVIまたはルーティングされたインターレースで、IP指向ブロードキャストとフォワードプロトコルを有効にする必要があります。
- B. クライアント範囲の IP ヘルパー アドレスは、WoL サーバーのサブネットが存在する SVI またはルーティング インターフェイスで有効にする必要があります。

- C. WoLサーバーサブネットが存在するSVIまたはルーティングインターフェースでは、クライアント範囲のIPヘルパーアドレスを無効にする必要があります。
- D. WoLサーバーのIPヘルパーアドレスは、クライアントサブネットが存在するSVIまたはルーティングインターフェースで有効にする必要があります。
- E. クライアントサブネットが存在するすべてのSVIまたはルーティングインターフェースで、IP指向ブロードキャストとフォワードプロトコルを無効にする必要があります。

正解: ([正解を表示します](#))

Wake-on-LAN across routed boundaries requires directed broadcast handling toward the client subnet and helper behavior to forward the magic packet correctly. A sleeping host normally does not maintain a normal IP session, so the WoL magic packet is commonly sent as a directed broadcast to the destination subnet. Cisco designs therefore enable IP directed broadcast on the routed interface or SVI for the client subnet where the sleeping endpoints reside, while controlling the feature carefully to avoid abuse. The ip forward-protocol behavior is also relevant because routers must forward the required UDP broadcast traffic rather than dropping it at the Layer 3 boundary. In addition, an IP helper address toward the WoL server or relay path may be needed on client-facing interfaces so the magic packet can be delivered from the management or server subnet to the client subnet. PortFast helps endpoints become reachable quickly at Layer 2, but it does not replace routed broadcast forwarding. Reference topics: Wake-on-LAN, directed broadcast, ip helper- address, UDP forwarding, PortFast, campus services.

質問: 114

ワイヤレスエンドポイントは、Cisco SD-AccessアーキテクチャのHTDBにどのように登録されますか？

- A. ファブリックエッジノードは、APからのCAPPWAPメッセージングに基づいてHTDBを更新します
- B. 新しいクライアントがワイヤレスネットワークに接続すると、ファブリックWLCがHTDBを更新します
- C. 境界ノードは最初にエンドポイントを登録し、次にHTDBを更新します
- D. ファブリックAPIは、クライアントのEIDおよびRLOCでHTDBを更新します

正解: ([正解を表示します](#))

Wireless endpoints in Cisco SD-Access are registered through the fabric wireless integration process, and the fabric WLC is the component that updates endpoint information as clients join the wireless network. The Host Tracking Database stores endpoint identifier and location information used by the SD-Access control plane.

For wired endpoints, the fabric edge detects the endpoint and registers the binding. For wireless endpoints, the fabric WLC is integrated with the fabric and has the client session context, so it provides the endpoint information needed for HTDB updates. Border nodes do not first register ordinary wireless endpoints; their role is to connect the fabric to external networks and other domains. Fabric APs also do not own the complete EID-to-RLOC registration function in the way the controller does. The design implication is that fabric wireless requires correct WLC integration with Cisco Catalyst Center, fabric site configuration, control-plane reachability, and policy integration through Cisco ISE. Client registration, mobility, and policy enforcement depend on this controller-to-fabric relationship being stable. Reference topics: SD-Access wireless, fabric WLC, Host Tracking Database, endpoint registration, LISP control plane.

質問: 115

Cisco SD-Accessネットワークファブリックのコントロールプレーンノードの目的は何ですか？

- A. エンドポイントデータベースとエンドポイントとエッジノード間のマッピングを維持するため
- B. ファブリック内のエンドポイントを検出し、ホスト追跡データベースにEIDからファブリックエッジへのノードバインディングを通知します
- C. ネットワークファブリック内のエンドポイントを識別および認証する
- D. ネットワークファブリックと外部ネットワーク間のネットワークゲートウェイとして機能する

正解: ([正解を表示します](#))

The SD-Access control-plane node maintains the endpoint database and resolves mappings between endpoints and the fabric edge nodes where those endpoints are located. Cisco SD-Access uses LISP in the fabric control plane. Edge nodes register endpoint identifiers with the control-plane node, and other fabric nodes query the control plane when they need to locate a destination endpoint. This control-plane function is commonly described through LISP Map-Server and Map-Resolver behavior. Option B is not correct because endpoint detection occurs at the edge node, not at the control-plane node. The edge is the device connected to users, access points, or endpoint-facing networks, so it learns local endpoint presence and registers that information.

Option C refers to authentication and identity, which is integrated with Cisco ISE. Option D describes the border-node role, because border nodes connect the fabric to external networks. The corrected answer is therefore A. A stable SD-Access design must provide redundant control-plane reachability and scale the control-plane role according to fabric size, mobility, and endpoint churn. Reference topics: SD-Access control plane, LISP Map-Server, LISP Map-Resolver, endpoint-to-location mapping.

質問: 116

ネットワーク エンジニアは、キャンパス OSPF 展開を最適化する必要があります。現在、タイプ 1 またはタイプ 2 LSA がエリア内で生成されるたびに、OSPF プロセスは SPT 全体を再計算する必要があります。どのソリューションが再計算プロセスを改善しますか。

- A. iSPF
- B. 素晴らしい
- C. SPF
- D. 中国

正解: A ([コメントを發表する](#))

Incremental SPF is the correct solution because the problem specifically states that Type 1 or Type 2 LSAs force the OSPF process to recompute the entire shortest path tree. Cisco OSPF documentation states that when changes to a Type 1 or Type 2 LSA occur, the entire SPT is normally recomputed, and that incremental SPF allows the system to recompute only the affected part of the tree. This improves convergence and reduces CPU use because the router does not have to run a full Dijkstra calculation for every topology change. BFD detects failures quickly, but it does not change how OSPF recalculates the SPT after the LSA arrives.

Standard SPF is the default full calculation behavior and therefore does not solve the optimization requirement. PRC is associated with partial route calculation behavior in some link-state contexts, but for this OSPF question the Cisco feature name is iSPF. Reference topics: OSPF incremental SPF, Type 1 LSA, Type 2 LSA, SPF optimization, campus OSPF scaling.

質問: 117

エンジニアはNETCONFおよびCisco NX-OSベースのデバイスを使用しています。エンジニアには、Cisco NX-OSにのみ関連する特定の機能をサポートするYANGモデルが必要です。エンジニアはどのモデルを選択する必要がありますか？

- A. Native
- B. IEEE
- C. OpenConfig
- D. IETF

正解: A ([コメントを發表する](#))

A Cisco native YANG model is the right choice when the engineer needs to configure a capability that is specific to Cisco NX-OS. YANG models fall broadly into open standards-based models and vendor-native models. IETF and OpenConfig models are useful for multivendor consistency, but they intentionally model common behavior and may not cover every platform-specific feature. Cisco native models are generated to represent Cisco operating system configuration and operational data more closely, including features that are unique to IOS XE, IOS XR, or NX-OS. Cisco developer guidance commonly positions native models as the model family to use when full platform feature coverage is required. IEEE is not the relevant model family for Cisco NX-OS feature-specific configuration in NETCONF. OpenConfig is useful when the design objective is common cross-vendor abstraction, but that is not the question's requirement. Therefore, for an NX-OS-only feature that must be configured through NETCONF/YANG, the engineer should choose the Cisco native model for NX-OS.

質問: 118

ある企業がサービスプロバイダと協力してBGPポリシーを設計しています。この企業はプロバイダとデュアルホーム接続しており、受信トラフィックがどのリンクを通過するかを制御したいと考えています。また、この企業はプロバイダに対して複数のネットワークをアドバタイズし、その伝播がそれ以上進まないようにする必要があります。これらの要件を満たすBGP属性はどれでしょうか？

- A. ASパス
- B. WITH
- C. コミュニティ
- D. ローカル設定

正解: C ([コメントを發表する](#))

BGP community is the correct attribute because it can support both policy requirements through coordination with the service provider. Communities are route tags that an upstream provider can match to apply routing policy, such as adjusting local preference to influence which dual-homed provider link is preferred for inbound traffic. Communities also include well-known values used to constrain route propagation. Cisco documents no-export as a community that prevents a route from being advertised outside the receiving AS, and no-advertise as a community that prevents advertisement to any peer. This directly matches the requirement that several customer-advertised networks must propagate no further. MED can influence inbound path selection in some dual-link designs, but it does not provide the propagation-control function. AS- path prepending can influence inbound traffic but is less precise and again does not solve the no-further- advertisement requirement. Local preference is normally applied inside an AS, not directly by the customer to external peers. Reference topics: BGP communities, no-export, no-advertise, provider policy, inbound traffic engineering.

質問: 119

エンジニアは、マルチベンダーネットワーク環境を管理するための標準主導のYANGモデルを探しています。

エンジニアはどのモデルを選択する必要がありますか？

A. ネイティブ

B. OpenConfig

C. IETF

D. IEEE NETCONF

正解: ([正解を表示します](#))

For a standards-driven YANG model in a multivendor environment, the engineer should choose IETF models.

IETF YANG models are produced through the Internet Engineering Task Force standards process and define vendor-neutral data structures for common network functions. This makes them the best answer when the requirement emphasizes standards-based operation across multiple vendors. OpenConfig models are also vendor-neutral and widely used for operational consistency, but OpenConfig is an operator-driven model set rather than an IETF standards body output. Cisco native models expose platform-specific Cisco IOS XE, NX- OS, or IOS XR features and often provide the most complete coverage for Cisco-specific capabilities, but they are not intended as the primary multivendor standard model. IEEE NETCONF is not the correct model category; NETCONF is a protocol, and IEEE is not the YANG model set normally selected for general device configuration in this context. In Cisco automation design, the model choice should match the use case: IETF for standards-driven multivendor workflows, OpenConfig for operator-neutral abstractions, and native models for vendor-specific feature completeness.

質問: 120

vEdgeルーターの冗長性が設計されている場合、どのFHRPがサポートされますか？

A. HSRP

B. OMP

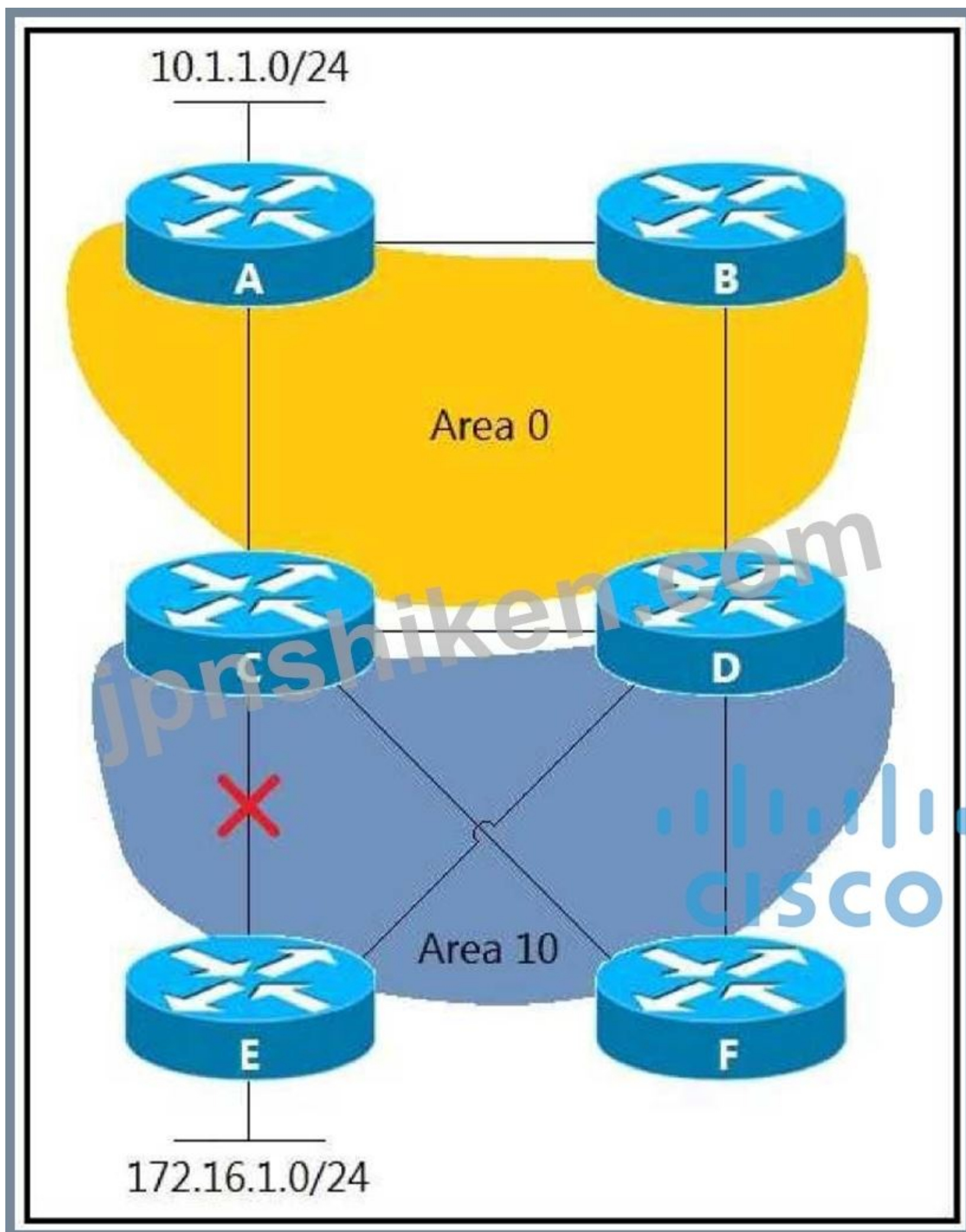
C. GLBP

D. VRRP

正解: ([正解を表示します](#))

When Cisco SD-WAN edge router redundancy is designed, VRRP is the supported first-hop redundancy protocol in the SD-WAN edge context. VRRP provides a virtual default gateway for LAN-side devices, allowing one WAN Edge router to act as the active gateway while another can take over if the active router or tracked condition fails. Cisco SD-WAN designs use VRRP with tracking and policy alignment so branch LAN traffic can fail over to the appropriate edge router while the SD-WAN overlay handles transport and tunnel selection. HSRP and GLBP are traditional campus first-hop redundancy protocols, but they are not the supported FHRP answer for vEdge router redundancy in this design context. OMP is the SD-WAN control- plane routing protocol used between WAN Edge routers and vSmart controllers; it is not a first-hop redundancy protocol for LAN hosts. Therefore, the correct FHRP for vEdge router redundancy is VRRP. The design should also ensure that VRRP priorities, tracking, and SD-WAN routing preferences align so traffic exits the intended edge during normal and failure conditions.

質問: 121



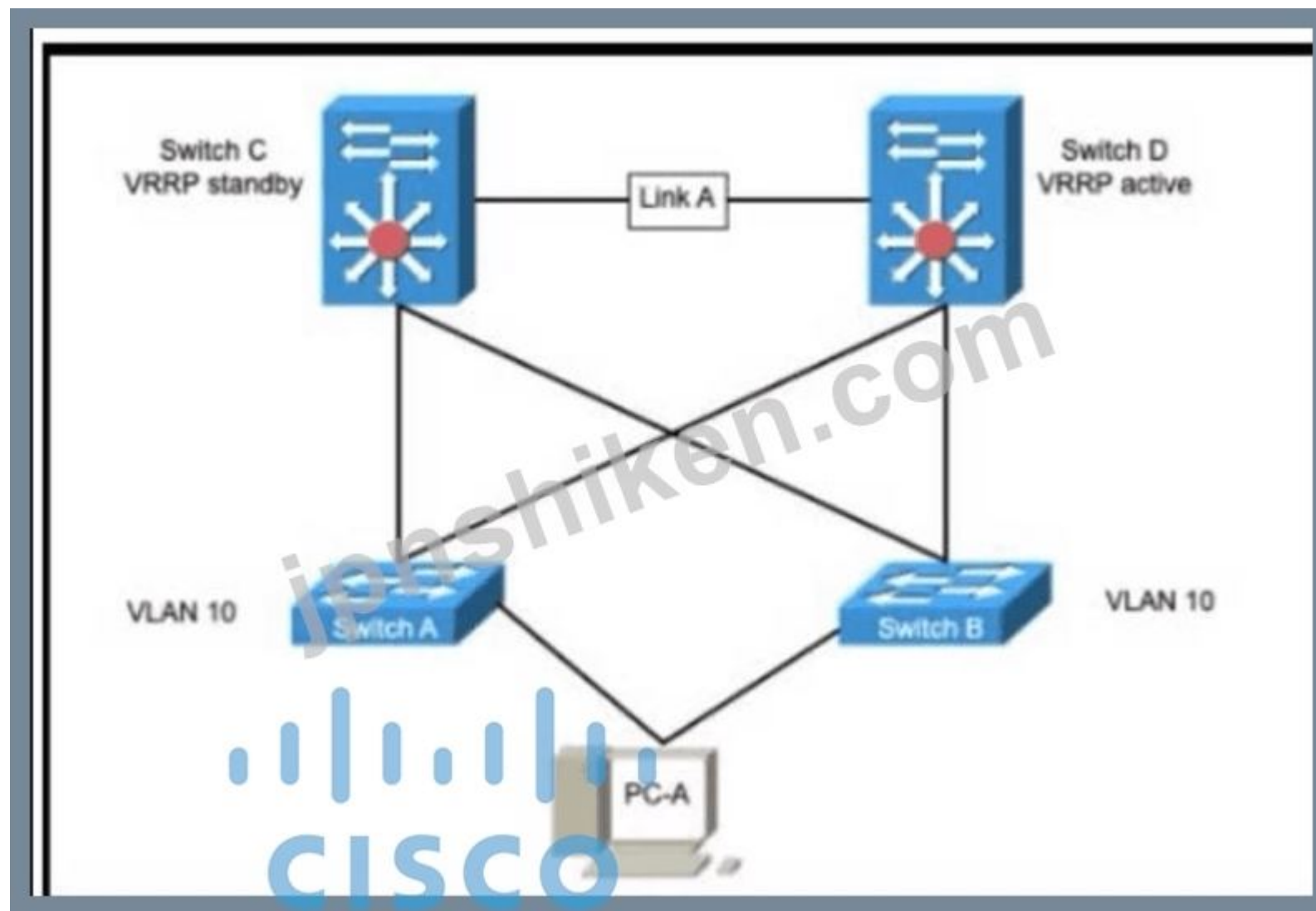
図を参照してください。エリア 10 は通常の OSPF エリアであり、ネットワーク 10.1.1.0/24 と 172.16.1.0/24 は内部です。
ルータ C と E 間のリンクに障害が発生した場合、両方のネットワーク間で最適なルーティングを提供する設計はどれですか。
A. ルータ C と D 間のリンクをエリア 10 に移動します。
B. ルータ E と F の間に OSPF 仮想リンクを作成します。
C. エリア 10 のルータ E と F の間にトンネルを作成します。
D. エリア 10 をそれほど短くないエリアにします。

正解: ([正解を表示します](#))

Moving the link between routers C and D into area 10 provides the most direct OSPF design correction for optimal routing between the internal networks after the C-E link fails. OSPF requires careful area design because intra-area routes are preferred over interarea routes, and traffic between two routers in the same regular area should not be forced through a less optimal backbone or interarea path when a better local path exists. If the C-D link remains outside area 10, the failure of C-E can cause traffic between 10.1.1.0/24 and 172.16.1.0/24 to take a suboptimal interarea route. Placing the C-D link in area 10 keeps the alternate path within the same area and allows OSPF to calculate the best intra-area route when the primary link fails. A virtual link is used to connect a disconnected area to area 0, not to optimize this intra-area failure case. A tunnel adds unnecessary overlay complexity, and making area 10 NSSA changes external LSA behavior, not internal path optimization. Reference topics: OSPF area design, intra-area preference, failure-path optimization, link placement.

有効的な**300-420J**問題集はJPNTest.com提供され、**300-420J**試験に合格することに役に立ちます！JPNTest.comは今最新**300-420J**試験問題集を提供します。JPNTest.com 300-420J試験問題集はもう更新されました。ここで**300-420J**問題集のテストエンジンを手に入れます。最新版のアクセス、<https://www.jpntest.com/shiken/300-420J-mondaishu> **457**問、**30%ディスカウント**、特別な割引コード：**JPNshiken**」

質問: 122



展示資料を参照してください。以下の要件を満たす必要があります。

* VLANは複数のアクセススイッチにまたがる。

* すべてのVLANは、すべてのアクセススイッチのアップリンク上で、ディストリビューションスイッチヘトランキングされます。

* STPバージョンはRapid PVST+です。

どの設計が最も速いスパニングツリー収束を実現しますか？

A. スイッチDはVLAN 10のセカンダリルートとして設定され、スイッチCはVLAN 10のプライマリルートとして設定され、リンクAはレイヤ2トランクとして設定されています。

B. スイッチDはVLAN 10のプライマリルートとして設定され、スイッチCはVLAN 10のセカンダリルートとして設定され、リンクAはレイヤ2トランクとして設定されています。

C. スイッチDはVLAN 10のプライマリルートとして設定され、スイッチCはVLAN 10のセカンダリルートとして設定され、リンクAはレイヤ3ルーティングリンクとして設定されています。

D. スイッチDはVLAN 10のセカンダリルートとして設定され、スイッチCはVLAN 10のプライマリルートとして設定され、リンクAはレイヤ3ルーティングリンクとして設定されています。

正解: [\(正解を表示します\)](#)

Switch D should be the VLAN 10 primary root, Switch C should be the VLAN 10 secondary root, and link A should remain a Layer 2 trunk. The requirements say VLANs span multiple access switches and all VLANs are trunked on all uplinks. In that design, the distribution interswitch link must remain Layer 2 so the VLAN can exist consistently across the access layer. For fastest and most deterministic Rapid PVST+ convergence, the spanning-tree root bridge should align with the active first-hop gateway for the VLAN. In the exhibit, Switch D is the active gateway for VLAN 10, so making Switch D the primary STP root avoids unnecessary traffic tromboning and keeps the forwarding path optimal. Switch C becomes the secondary root for resiliency. A Layer 3 routed interdistribution link is a best practice when VLANs are not stretched, but it contradicts the stated requirement that VLANs span access switches. Reference topics: Rapid PVST+, STP root placement, FHRP alignment, Layer 2 campus distribution, VLAN spanning design.

エンジニアは、定義済みのビジネス要件に基づいてトラフィックをクラスに分離できるスケーラブルな QoS アーキテクチャを設計する必要があります。また、設計では、差別化サービス コード ポイントを QoS 優先度記述子値として使用し、少なくとも 10 レベルの分類をサポートする必要があります。エンジニアはどの QoS テクノロジを設計に含める必要がありますか？

- A. 返信
- B. ディフサーブ
- C. ベストエフォート
- D. インターサーブ

正解: ([正解を表示します](#))

DiffServ is the correct scalable QoS architecture for classifying traffic according to business requirements and using DSCP values as the priority descriptor. Cisco QoS design distinguishes DiffServ from IntServ. IntServ uses per-flow signaling with RSVP and can reserve resources, but that model is not operationally scalable across large enterprise networks. DiffServ is class-based. Traffic is classified and marked at the network edge, then each device applies per-hop behavior to those classes using local QoS policy. DSCP provides a six-bit field in the IP header, which supports more than ten classification values and allows standardized treatment classes such as EF, AF, class selector, and default. Best effort provides no differentiated treatment, and RSVP is a signaling mechanism associated with IntServ rather than the requested scalable class-based architecture.

"Interserv" is not the correct Cisco QoS model term. Therefore, the engineer should design around DiffServ, with a clear trust boundary, marking policy, queuing model, and consistent treatment across WAN and campus boundaries. Reference topics: DiffServ, DSCP, per-hop behavior, class-based QoS, scalable QoS architecture.

質問: 124

エンジニアは、複数のレイヤー3仮想ネットワークを作成する機能を含むキャンパスネットワークのソリューションを提案する必要があります。各ネットワークには、データ転送用の独自のアドレス指定構造とルーティングテーブルが必要です。このソリューションは、数百の仮想ネットワークをサポートし、最小限の管理オーバーヘッドで簡単な構成と管理を可能にするようにスケーラブルである必要があります。エンジニアはどのソリューションを推奨しますか？

- A. ホップバイホップEVN
- B. マルチホップMPLSコア
- C. マルチホップIPsecトンネリング
- D. ホップバイホップVRF-Lite

正解: A ([コメントを发表する](#))

Easy Virtual Network is the best fit when a campus needs many Layer 3 virtual networks with separate addressing and routing tables but wants lower operational overhead than manually configuring VRF-Lite everywhere. EVN builds on VRF concepts and simplifies virtualization across a hop-by-hop campus design by reducing repetitive configuration and enabling consistent virtual network transport across participating devices. Each virtual network maintains its own routing context, which supports overlapping addressing and segmentation. Hop-by-hop VRF-Lite can provide similar forwarding separation, but it becomes operationally heavy as the number of VRFs grows because each trunk, routing process, and interface association must be configured and maintained manually. Multihop MPLS can scale, but it introduces provider-style complexity and may require features or skills beyond a typical campus requirement. Multihop IPsec tunneling is not an efficient campus segmentation solution for hundreds of virtual networks. Therefore, EVN is the recommended design for scalable Layer 3 virtualization with simpler configuration and management. Reference topics: campus virtualization, Easy Virtual Network, VRF separation, scalable virtual networks, route isolation.

質問: 125

すべてのサイトにインターネットへの直接アクセスが実装されたSD-WANネットワークファブリックを設計する2つの利点は何ですか？ 2つ選択してください。）

- A. パブリッククラウドサービスプロバイダーがホストするアプリケーションのレイテンシを低減します。
- B. インターネット回線の待ち時間を短縮します。
- C. ゼロタッチプロビジョニングにより、サイト展開の配信速度を向上させます。
- D. インターネット回線で利用可能な帯域幅の合計を増やします。
- E. MPLS回線上のネットワークトラフィックを軽減します。

正解: ([正解を表示します](#))

Implementing direct Internet access at every Cisco SD-WAN site improves application experience and reduces unnecessary backhaul. Many modern business applications are hosted in SaaS or public cloud environments. If branch Internet traffic must hairpin through a central data center before reaching those applications, latency increases and central WAN or security stacks can become bottlenecks. SD-WAN DIA allows selected cloud and Internet-bound traffic to exit locally while still applying policy, segmentation, and security controls. This decreases latency to applications hosted by public cloud providers and improves user experience. It also alleviates traffic on MPLS circuits because Internet and cloud traffic no longer needs to consume private WAN bandwidth back to a centralized egress point. DIA does not inherently reduce the physical latency of an Internet circuit, and it does not increase the provisioned circuit bandwidth. Zero-touch provisioning may accelerate site deployment, but that is a separate SD-WAN onboarding benefit rather than a specific benefit of deploying DIA at every site. Therefore, the two correct benefits are reduced latency to public cloud applications and reduced load on MPLS circuits.

質問: 126

gRPC はどのような統合機能を提供しますか？

- A. 認証とディレクトリ サービスに LDAP プロトコルを活用し、RPC 通信における安全なアクセス制御を確保する
- B. クライアントとサーバーアプリケーション間のリアルタイムメッセージングとコラボレーションにXMPPプロトコルを活用する
- C. プロトコルバッファを活用して、ネットワーク上で構造化データの効率的なシリアル化と逆シリアル化を実現する
- D. GRAPH-API を活用してネットワークの監視と管理を行い、RPC 関連のメトリックとパフォーマンス統計を包括的に可視化します。

正解: C ([コメントを発表する](#))

gRPC provides efficient integration by using protocol buffers to serialize and deserialize structured data over the network. In network programmability and telemetry designs, gRPC is commonly used as the transport framework for streaming model-driven telemetry or for remote procedure calls between collectors, controllers, and network devices. Protocol buffers define compact structured messages that are faster and less ambiguous than parsing text output. This makes gRPC suitable for high-volume telemetry and programmatic device interaction. LDAP is a directory and authentication-related protocol, not the integration mechanism provided by gRPC. XMPP is associated with messaging and presence systems, not gRPC data serialization.

Graph API is not a gRPC feature. The operational value of gRPC is that it provides a modern RPC framework with efficient message encoding, service definitions, and support for secure transport through TLS when implemented by the platform. Reference topics: gRPC, protocol buffers, model-driven telemetry, network programmability, structured data serialization. It improves controller interoperability.

質問: 127

SO-WAN アーキテクチャで HOC 接続を最小限に抑え、vSmart コントローラの負担を軽減する機能はどれですか？

- A. 制御接続
- B. 配線エラー
- C. 色
- D. 親和性

正解: D ([コメントを発表する](#))

Affinity is the SD-WAN feature used to control and reduce control connections between WAN Edge devices and vSmart controllers. Cisco SD-WAN high-availability documentation describes controller affinity as a way to limit the number of controllers that an edge device establishes control connections with, while still preserving controller and data center redundancy. This reduces unnecessary control-plane load and limits the strain placed on vSmart controllers in large overlays. The feature is particularly useful when multiple vSmart controllers exist across multiple data centers and the designer wants predictable primary and backup controller relationships. The color attribute identifies transport characteristics, such as public Internet, MPLS, or other transport types. Control-connections is a general concept, not the feature that minimizes them. Control- direction is not the correct design construct for this purpose. A scalable SD-WAN design should use affinity where controller scale, regionalization, and control-plane efficiency must be explicitly managed. Reference topics: Cisco SD-WAN controller affinity, vSmart scale, control connections, TLOC control-plane optimization.

質問: 128

アーキテクトは、次の要件に基づいて 4 つのサブネット間で 40 Mbps のインターネット接続が共有されるように、顧客向けの QoS ソリューションを作成する必要があります。

* 各サブネットは、トラフィックのピーク時に 10 Mbps 以上のダウンロード帯域幅を受信する必要があります。

* 他のサブネットがアイドル状態の場合、サブネットはトラフィックのピーク時以外に最大 40 Mbps を使用できます。

* ダウンロード トラフィックに遅延が発生してはなりません。
アーキテクトはどのソリューションを選択する必要がありますか？

- A. レート制限とシェーピング
- B. 帯域幅のパーセンテージとポリシング
- C. シェーピングとポリシング
- D. 帯域幅のパーセンテージとレート制限

正解: ([正解を表示します](#))

Bandwidth percentage with policing is the correct QoS approach. The requirement has two separate parts:

each subnet must receive a minimum share during congestion, and no subnet should introduce delay into download traffic. Bandwidth percentage is used to allocate a minimum bandwidth guarantee among classes during congestion while still allowing a class to use unused bandwidth when other classes are idle. Policing enforces a rate by dropping or remarking excess traffic rather than buffering it. That behavior is important because the question explicitly says download traffic must never experience delay. Shaping would be wrong because shaping buffers excess packets and transmits them later, which intentionally adds delay. Rate-limiting is a broad term, but when the design choice offers policing explicitly, policing is the precise mechanism that matches the no-delay requirement. A parent or class policy can reserve 25 percent per subnet on a 40 Mbps service, giving each subnet 10 Mbps during congestion while unused bandwidth remains available. Reference topics: CBWFQ bandwidth percent, traffic policing, shaping versus policing, Internet edge QoS, congestion management.

質問: 129

エンジニアは、ブランチ サイトを持つ顧客向けにインバンド管理ソリューションを設計する必要があります。ソリューションでは、MPLS WAN 上の管理プロトコルを使用してブランチ サイトのリモート管理を可能にする必要があります。キューイングは、次のクラスを使用してリモート サイトで実装されます。

- Class1 equals voice traffic
- Class2 equals mission-critical traffic
- Class3 equals default traffic

ソリューションでは、WAN 上の管理トラフィックをどのように優先順位付けする必要がありますか？

- A. トラフィックを DSCP CS1 でマークし、Class3 で使用可能な帯域幅を削減して割り当てられた最小帯域幅で Class2 にマップします。
- B. トラフィックをDSCP CS6でマークし、クラス2に利用可能な帯域幅を減らして割り当てられた最小帯域幅でクラス1にマップします。
- C. トラフィックを DSCP EF でマークし、Class2 で使用可能な帯域幅を削減して割り当てられた最小帯域幅で Class1 にマップします。
- D. トラフィックをDSCP CS2でマークし、クラス3に利用可能な帯域幅を減らして割り当てられた最小帯域幅でクラス2にマップします。

正解: ([正解を表示します](#))

For in-band management over the production WAN, the proper marking is DSCP CS2 mapped into the management or operations class with a minimum bandwidth guarantee. Cisco enterprise QoS guidance separates network control traffic from network management traffic. Routing protocol and control-plane traffic should not be diluted by ordinary management flows. CS6 is typically reserved for routing and network-control traffic; EF is reserved for low-latency voice bearer traffic; CS1 is commonly used for scavenger or low-priority traffic. Network management flows such as SSH, SNMP, NTP, FTP-based maintenance, and syslog require reliable delivery, but they should not starve business traffic or real-time services. Therefore, marking with CS2 and mapping the traffic into Class2 is the correct design choice. The bandwidth adjustment should come from a lower-priority class, not from the routing/control or real-time class. Because this is in-band management, the architect must also ensure access control, source restrictions, AAA, logging, and management-plane protection. QoS classification alone does not secure management access; it only gives the traffic an appropriate forwarding treatment over the MPLS WAN. Reference topics: in-band management, DSCP CS2, enterprise QoS classes, management-plane traffic.

質問: 130

DMVPNトンネルフラップが断続的に発生する一般的な問題はどれですか。

- A. ルーティングネイバーの到達可能性の問題
- B. 次善のルーティングテーブル
- C. インターフェース帯域幅の輻輳
- D. ハブルーターへのGREトンネルは暗号化されていません

正解: **A** ([コメントを發表する](#))

Intermittent DMVPN tunnel flaps are commonly caused by reachability problems between routing neighbors or between the tunnel endpoints that support those neighbors. DMVPN depends on several layers working correctly: the underlay transport, GRE tunnel interface, NHRP registration and resolution, IPsec protection if used, and the routing protocol running across the tunnel. If the routing neighbor cannot remain reachable, the tunnel may appear to flap even though the encryption profile itself is not the root cause. A suboptimal routing table can produce poor forwarding or asymmetric traffic, but it is not the most common direct cause of repeated tunnel up/down events. Bandwidth congestion can degrade performance, but a correctly designed tunnel does not flap merely because utilization is high unless keepalives, NHRP, or routing packets are being lost. The fact that a GRE tunnel is not encrypted is a security design problem, not a tunnel-flap cause. The right troubleshooting posture is to verify stable underlay reachability, tunnel interface state, NHRP mappings, and routing adjacency stability before assuming an IPsec encryption failure.

質問: **131**

Cisco SD-WAN ネットワーク ファブリックでの TLOC 拡張の目的は何ですか？

- A.** サイト内の WAN エッジ ルーターの冗長性を促進するため
- B.** WAN エッジ ルーターが WAN トランスポート ネットワークに接続する物理インターフェイスを特定します。
- C.** ネットワーク トランスポート インターフェイスに適用される可能性のある色の数を拡張します。
- D.** 複数の物理インターフェイスを単一の論理インターフェイスに集約する

正解: ([正解を表示します](#))

A TLOC extension facilitates WAN Edge router redundancy within a site. In Cisco SD-WAN, a Transport Location identifies a WAN transport attachment using attributes such as system IP, color, and encapsulation.

A TLOC extension allows one WAN Edge router to use a transport circuit physically connected to another WAN Edge router at the same site. This is valuable when a branch has two edge routers and two transports, but each transport is cabled to only one router. With TLOC extension, each router can still gain logical access to both transports, improving redundancy and allowing better path diversity without adding more service- provider circuits or handoffs. The feature does not simply identify the physical interface; that is part of TLOC definition. It does not increase the number of colors or create a port-channel-style aggregate. A sound design also accounts for failure domains, service-side routing, and whether the inter-router link has enough bandwidth for extended transport traffic. Reference topics: Cisco SD-WAN TLOC, TLOC extension, WAN Edge redundancy, transport color, branch resiliency.

質問: **132**

ISPは、MPLSを介したレイヤ3 VPNサービスを、4つのブランチと各ブランチに複数のCEルータを持つ顧客に提供します。CEルータから学習したルートを交換するには、ISPがPEルータ間でアクティブ化する必要があるBGPアドレスファミリーはどれですか。

- A.** アドレスファミリーマルチキャスト
- B.** L2VPN EVPN
- C.** VPNv4ユニキャスト
- D.** IPv4ユニキャスト

正解: ([正解を表示します](#))

The provider must activate the VPNv4 unicast address family between PE routers. In an MPLS Layer 3 VPN, customer routes learned from CE routers are placed into VRFs on the PE routers. To carry those customer routes across the provider backbone without overlap, MP-BGP advertises VPNv4 routes, which combine an IPv4 prefix with a route distinguisher. Route targets then control which receiving VRFs import the routes.

Plain IPv4 unicast would not preserve customer VPN context and would not solve overlapping route separation. L2VPN EVPN is used for Ethernet VPN services and data center or Layer 2 VPN use cases, not classic MPLS L3VPN route exchange between PEs. Address-family multicast is unrelated to exchanging customer unicast routes between PE routers. Cisco MPLS VPN design relies on MP-BGP VPNv4 between provider edge routers, while PE-CE routing can be static, OSPF, EIGRP, BGP, or another supported method.

Therefore, VPNv4 unicast is the correct PE-to-PE address family. Reference topics: MPLS Layer 3 VPN, MP- BGP, VPNv4 unicast, route distinguishers, route targets, PE-CE routing.

質問: **133**

Cisco Catalyst SD-WAN Cloud OnRampの目的は何ですか？

- A. あらゆる種類のリンクやプロバイダーに依存しない、安全でプライベートな接続を提供します。
- B. Cisco Catalyst SD-WANのアプリケーション対応ポリシーとルートクリティカルアプリケーションを作成します。
- C. マルチクラウド環境全体で、ビジネスに不可欠なアプリケーションへの安全なアクセスを提供します。
- D. オンプレミス環境をクラウドに接続するプロセスを簡素化し、自動化します。

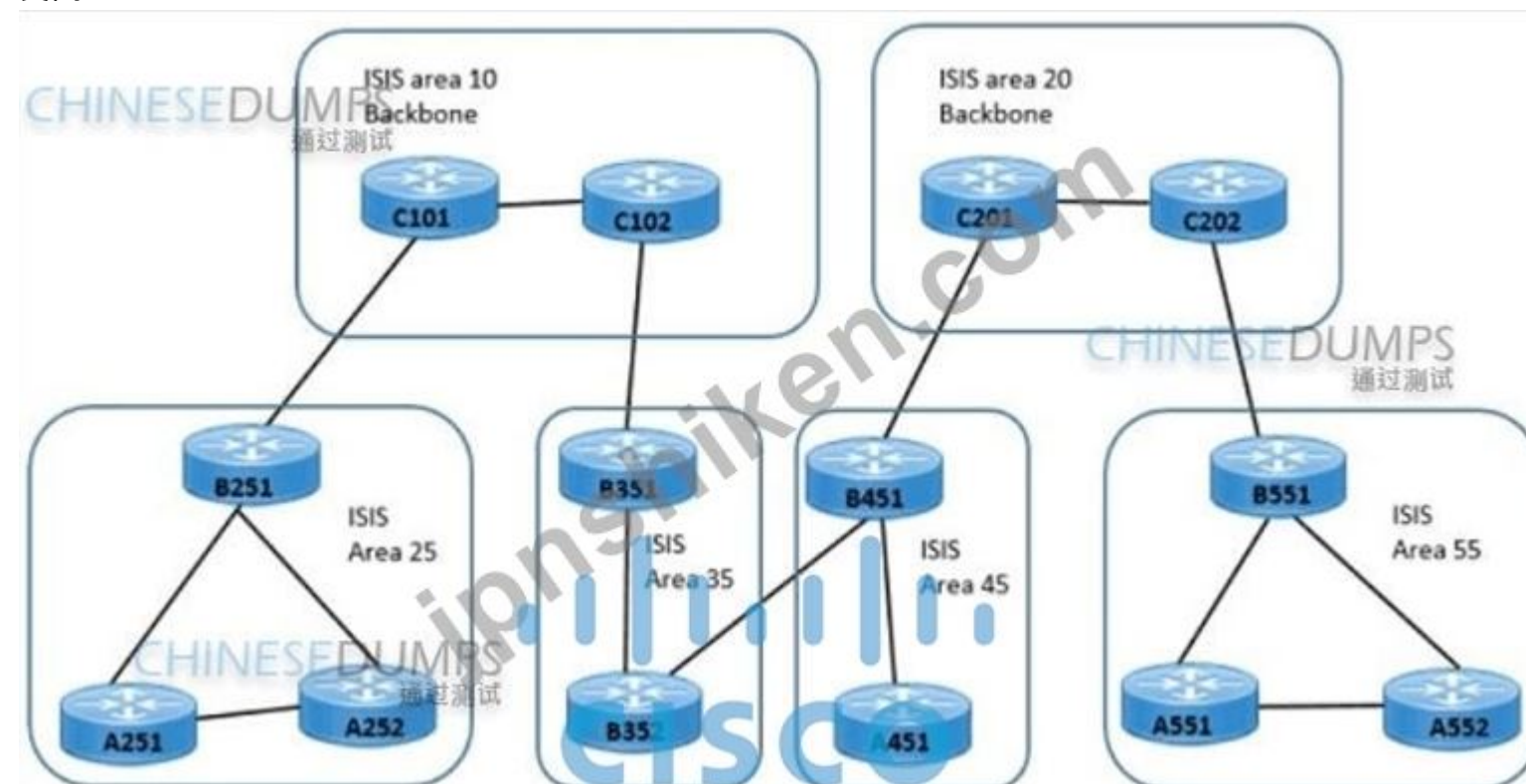
正解: (正解を表示します)

Cisco Catalyst SD-WAN Cloud OnRamp is built to simplify and automate cloud connectivity from enterprise sites. Cisco describes the feature as a way to simplify connecting on-premises environments to cloud-hosted applications and workloads while extending SD-WAN policy into the cloud. The key design value is automation: branches and data centers can connect to SaaS, IaaS, or multicloud resources without manually building every tunnel, route policy, and cloud attachment. That is why option D is the strongest answer.

Option A is too broad because generic transport independence is a core SD-WAN value, not the specific purpose of Cloud OnRamp. Option B confuses Cloud OnRamp with application-aware routing policy.

Option C describes a business outcome, but the operational purpose tested here is automated on-premises-to-cloud integration and consistent policy extension. Reference topics: Cisco Catalyst SD-WAN Cloud OnRamp, multicloud connectivity, cloud workload access, SD-WAN fabric extension, automated cloud onboarding.

質問: 134



展示品を参照してください。エンジニアは、次のような要件を持つ企業顧客向けに階層型ISISソリューションを設計しています。

- * エリア25と55のユーザーは、両方のバックボーンエリアからトラフィックを送受信します。
- * エリア35と45のリンクフラップは他のエリアに影響を与えてはならない
- * エリア35と45では、今後12か月以内にルーターが2倍になる。

エンジニアはどの設計を選択する必要がありますか？

- A. A シリーズ ルータ レベル 2、B シリーズ ルータ レベル 2、C シリーズ ルータ レベル 1
- B. A シリーズ ルータ レベル 1/2、B シリーズ ルータ レベル 2、C シリーズ ルータ レベル 2
- C. A シリーズ ルータ レベル 1。B シリーズ ルータ レベル 1/2。C シリーズ ルータ レベル 2
- D. A シリーズ ルータ レベル 1.2、B シリーズ ルータ レベル 1/2、C シリーズ ルータ レベル 1/2

正解: (正解を表示します)

The IS-IS hierarchy should place access or internal area routers at Level 1, boundary routers at Level 1/2, and backbone routers at Level 2. In this exhibit, users in areas 25 and 55 must send and receive traffic through both backbone areas, while link flaps in areas 35 and 45 must not affect other areas. Cisco IS-IS design uses Level 1 areas to contain local topology changes and Level 2 routing to interconnect areas through the backbone. Level 1/2 routers sit at the area boundary and leak or advertise reachability between the local Level 1 area and the Level 2 backbone. Option C matches that hierarchy by placing A-series routers as Level 1, B-series routers as Level 1/2, and C-series routers as Level 2. Making too many routers Level 1/2 would expand the impact of topology changes and reduce scaling benefits. Making backbone routers Level 1 would break interarea design. Reference topics: IS-IS hierarchy, Level 1, Level 2, Level 1/2 routers, failure-domain containment, enterprise scaling.

質問: 135

```
bandwidth percent 26
random-detect dscp-based
class BULK
bandwidth percent 5
random-detect dscp-based
class SCAVENGER
bandwidth percent 1
class class-default
bandwidth percent 24
random-detect
!
class-map match-all BULK
match ip dscp af11 af12
class-map match-all VIDEO
match ip dscp af41 af42
class-map match-any ROUTING
match ip dscp cs6
class-map match-all MISSION-CRITICAL
match ip dscp af21 af22
class-map match-any SIGNALLING
match ip dscp cs3
match ip dscp af31
class-map match-all VOICE
match ip dscp ef
class-map match-all SCAVENGER
match ip dscp cs1
!
interface GigabitEthernet0/2
description Link to DC
service-policy output WAN-DC-LINK
```

図を参照してください。顧客は、GigabitEthernet0/2 インターフェイスを通過するネットワーク管理トラフィックに QoS を適用する必要があります。8 つのキューイング クラスがすべて使用されているため、新しい要件を既存のポリシーに統合する必要があります。顧客はどのソリューションを選択する必要がありますか？

- A. トラフィックを DSCP CS5 にマークし、SIGNALLING クラスに割り当てます。次に、既存のキュー サイズを基準にして、SIGNALLING クラスに追加の帯域幅をプロビジョニングできるかどうかを判断します。
- B. トラフィックを DSCP CS4 にマークし、SIGNALLING クラスに割り当てます。次に、クラス内でトラフィックに優先順位を付けます。
- C. トラフィックを DSCP CS6 にマークし、ROUTING クラスに割り当てます。次に、クラス内でトラフィックに優先順位を付けます。
- D. トラフィックを DSCP CS2 にマークし、ROUTING クラスに割り当てます。次に、既存のキュー サイズを基準にして、ROUTING クラスに追加の帯域幅をプロビジョニングできるかどうかを判断します。

正解: [\(正解を表示します\)](#)

Network management traffic should be marked and treated as a dedicated operations class, not mixed into voice signaling or real-time queues. Cisco enterprise QoS design guidance commonly reserves DSCP CS2 for operations, administration, and management traffic such as SNMP, SSH, syslog, NTP, and other management flows. The question states that all eight queuing classes are already used, so the design should integrate management traffic into the closest existing class and validate that the queue has enough bandwidth. Option D is the best fit because it uses CS2 and requires baselining before allocating more bandwidth. CS6 is normally used for network control and routing protocols; using it for ordinary management traffic would compete with routing adjacencies and control-plane stability. CS5 and CS4 are also unsuitable because they are normally associated with video/signaling-related classes in many enterprise QoS models. The professional design action is to classify management traffic explicitly, mark it CS2 at the trust boundary, map it to the existing class, and then measure queue utilization before changing bandwidth allocations. Reference topics: Cisco QoS baseline, DSCP CS2, network management traffic, queue bandwidth design.

質問: 136

モデル駆動型テレメトリの考慮事項を左側から右側の適用先のモードにドラッグ アンド ドロップします。

uses a transient connection

no need to open ports for inbound management traffic

anycast and load-balancing

single channel (config and streaming)

Dial-In Mode

Dial-Out Mode

正解:

uses a transient connection

no need to open ports for inbound management traffic

anycast and load-balancing

single channel (config and streaming)

Dial-In Mode

uses a transient connection

single channel (config and streaming)

Dial-Out Mode

anycast and load-balancing

no need to open ports for inbound management traffic

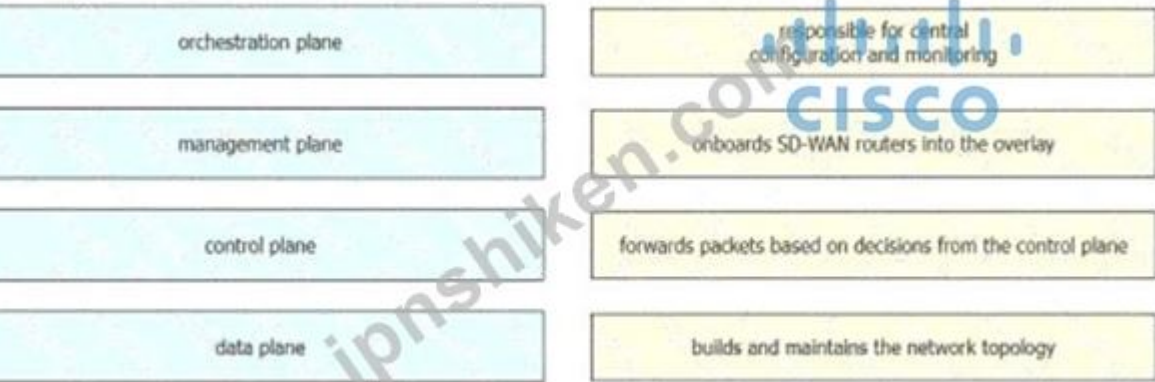
Explanation:

The telemetry mapping separates dial-in and dial-out behavior. In dial-in telemetry, the collector initiates the session toward the network device and subscribes to the data that it wants streamed. This model gives the collector direct control over the subscription lifecycle, but it requires the collector to reach the device and manage each session. In dial-out telemetry, the network device initiates the session toward the collector based on a configured subscription. This is useful when devices must stream telemetry to a centralized destination through firewalls or controlled management paths. Cisco model-driven telemetry can stream structured data from YANG models using subscription behavior rather than relying on repeated CLI polling. The selected drag-and-drop answer matches each consideration to the mode that owns the connection initiation and subscription handling. In design work, the architect must consider reachability direction, security policy, NAT traversal, collector scale, subscription persistence, and whether subscriptions should be configured on the device or created by the collector. Reference topics: model-driven telemetry, dial-in mode, dial-out mode, YANG subscriptions, gRPC telemetry, operational data streaming.

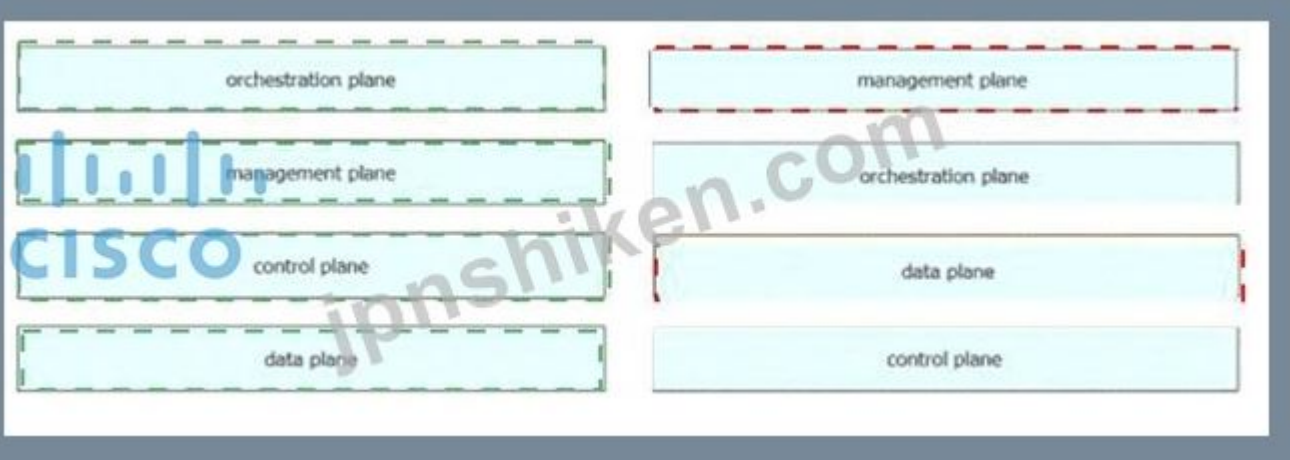
有効的な**300-420J**問題集はJPNTest.com提供され、**300-420J**試験に合格することに役に立ちます！JPNTest.comは今最新**300-420J**試験問題集を提供します。JPNTest.com 300-420J試験問題集はもう更新されました。ここで**300-420J**問題集のテストエンジンを手に入れます。最新版のアクセス、<https://www.jpntest.com/shiken/300-420J-mondaishu> **457**問、**30%ディスカウント**、特別な割引コード：**JPNshiken**」

質問: 137

左側の Cisco SD-WAN コンポーネントを右側の定義にドラッグ アンド ドロップします。



正解:



Explanation:

Cisco SD-WAN components map to distinct control, management, orchestration, and forwarding roles.

vManage is the centralized management plane used for configuration templates, monitoring, software management, and operational visibility. vSmart is the centralized control-plane component that distributes OMP routes, TLOC information, security parameters, and centralized policy to WAN Edge devices. vBond is the orchestration component used during onboarding; it authenticates devices, helps them discover controllers, and assists with NAT traversal. The WAN Edge router, whether physical or virtual, forms the data plane at the branch, campus, or data center and forwards user traffic through secure overlay tunnels. Correctly identifying these roles is essential because SD-WAN scale and resiliency depend on separating management, control, orchestration, and forwarding responsibilities. A design that confuses vManage with vSmart or vBond will misplace policy, onboarding, and forwarding functions. Reference topics: Cisco SD-WAN architecture, vManage, vSmart, vBond, WAN Edge, OMP, overlay control plane.

質問: 138

お客様は、ネットワークコンサルタントとQoS要件について話し合っています。お客様は、エンドツーエンドのパス検証が必須であることを指定しています。この要件を満たすQoSソリューションはどれですか。

- A. IntServ model with RSVP to support the traffic flows
- B. DiffServ model with PHB to support the traffic flows

C. marking traffic at the access layer with DSCP to support the traffic flows

D. marking traffic at the access layer with CoS to support the traffic flows

正解: [\(正解を表示します\)](#)

The requirement for end-to-end path verification identifies the Integrated Services model with RSVP. IntServ uses per-flow resource reservation, and RSVP signals the path between sender and receiver so devices along that path can admit or reject the requested service based on available resources. This makes IntServ appropriate when the design requires explicit end-to-end path validation and reservation for application flows.

DiffServ works differently. In a DiffServ design, traffic is marked, normally with DSCP or CoS at the edge, and each hop applies a configured per-hop behavior. DiffServ is scalable and common in enterprise WAN and campus QoS designs, but it does not perform per-flow signaling or verify that every device in the path can reserve the requested resources. Marking traffic at the access layer is an important QoS design practice, but marking alone is not path verification. Therefore, when the customer explicitly requires end-to-end path verification for business-critical traffic, Cisco QoS architecture points to IntServ with RSVP rather than a pure DiffServ marking model.

質問: **139**

エンジニアは、レイヤ3ルータが1つしかない小さなブランチサイト用にEIGRPネットワークを設計しています。エンジニアは、ルータがローカルLAN上で不要なマルチキャストメッセージを送信することなく、リモートEIGRPネイバーにローカルLANネットワークをアドバタイズすることを望んでいます。エンジニアはどのような行動を取るべきですか？

A. EIGRPの代わりに、このサイトに静的なデフォルトルートを使用します

B. networkコマンドとパッシブインターフェイス機能を使用してローカルLANをアドバタイズします

C. redistributeconnectedコマンドを使用してローカルLANネットワークを再配布します

D. ローカルLANサブネットをスタブネットワークとしてアドバタイズします

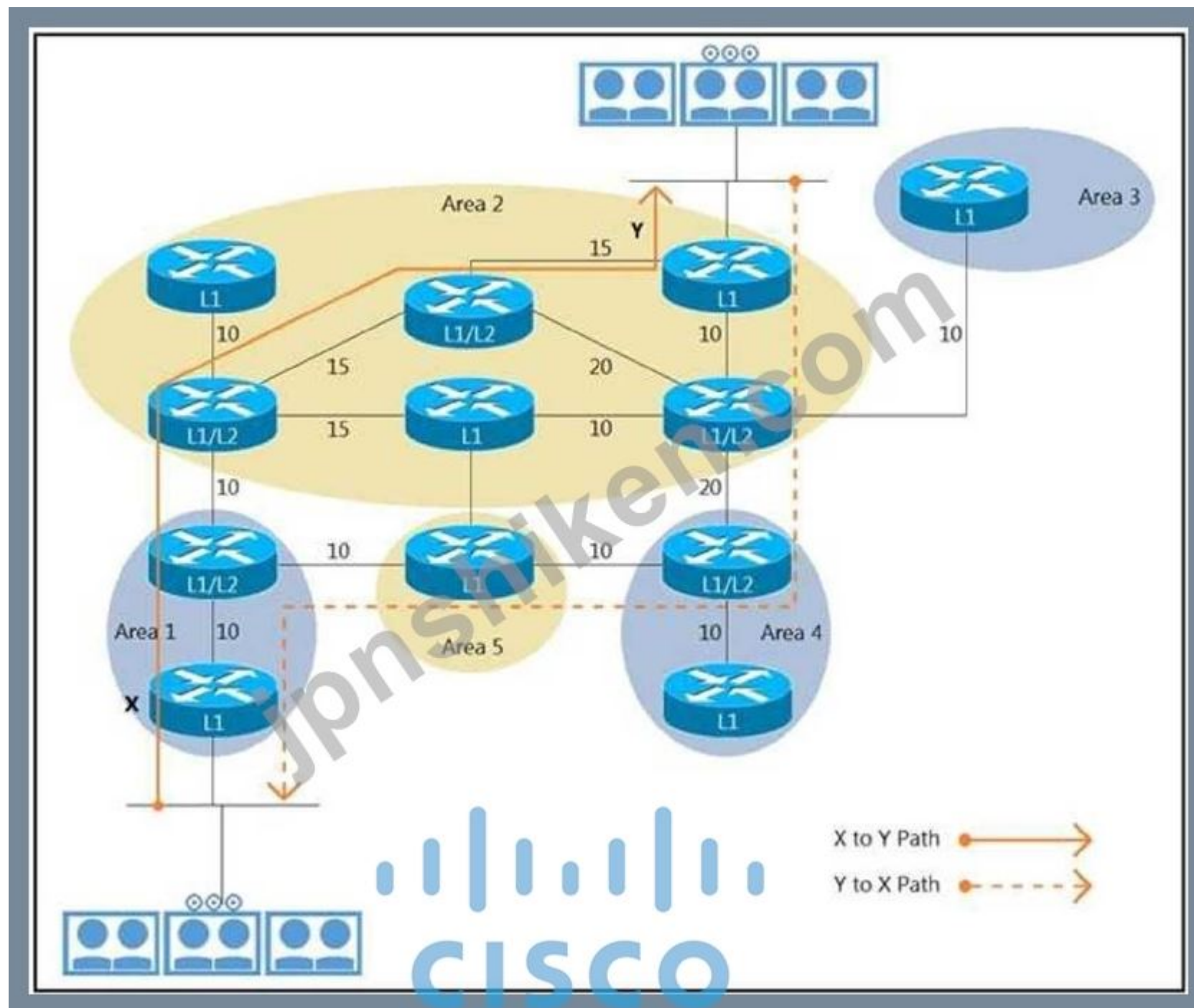
正解: **B** ([コメントを發表する](#))

The branch router should advertise the local LAN with the EIGRP network command and make the LAN- facing interface passive. A passive interface allows the connected network to be included in EIGRP advertisements while suppressing EIGRP hello packets and neighbor formation on that interface. This is exactly what the engineer needs: remote EIGRP neighbors learn the LAN subnet, but unnecessary multicast EIGRP messages are not sent onto the user LAN where no EIGRP neighbors should exist. Replacing EIGRP with a static default route would not meet the requirement to advertise the local LAN through EIGRP.

Redistributing connected routes can work, but it is less clean for a simple LAN interface and can introduce external-route behavior or policy complexity. Advertising the subnet as a stub network does not suppress multicast hellos on the LAN interface by itself. Cisco design best practice is to use passive interfaces on LAN- facing interfaces that should be advertised but should not form routing adjacencies.

Reference topics: EIGRP passive interface, network statements, branch routing, unnecessary neighbor prevention, multicast control traffic.

質問: **140**



展示を参照してください。顧客から、2つの場所間でポイントツーポイントのテレプレゼンスビデオ通話を行うと、ビデオ品質が低下し、遅延が発生するという報告があります。設計者は、トラフィックが出口トラフィックフローと入口トラフィックフローで同じパスをたどるように設計を最適化する必要があります。どの手法で設計を最適化しますか。

- A. エリア 2 のルーターでルートリークを設定します。
- B. エリア 1 のルーターでルートリークを設定します。
- C. エリア 4 のルーターに高メトリックを設定します。
- D. エリア 4 のルーターにルートフィルターを設定します。

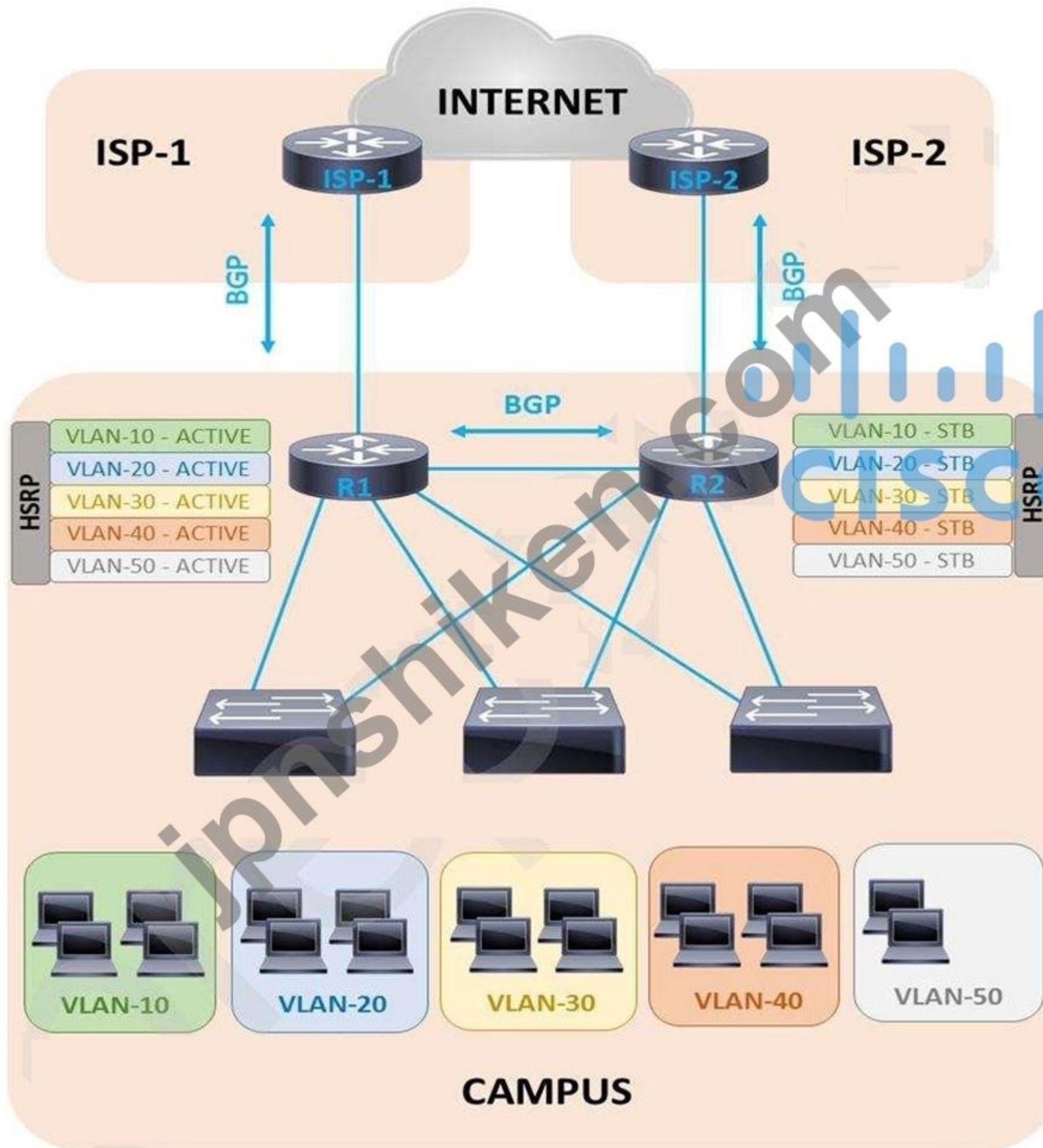
正解: **C** ([コメントを发表する](#))

For point-to-point telepresence, asymmetric routing can create delay variation, packet reordering, and troubleshooting complexity. The design requirement is to make egress and ingress flows follow the same path. In the topology, the best way to steer traffic away from the less desirable return path is to configure a higher metric on the router in area 4 so OSPF path selection favors the symmetric path. OSPF chooses paths based on cumulative cost, and increasing the cost on the unwanted path makes the preferred path more attractive without removing redundancy. Route leaking on the routers in area 1 or area 2 would introduce additional interarea route visibility but would not necessarily solve the asymmetric flow problem. A route filter in area 4 could remove reachability or create suboptimal failure

behavior rather than simply preferring the right path. Raising the metric is a cleaner design technique because it preserves backup reachability while influencing normal forwarding. Reference topics: OSPF cost design, symmetric routing, telepresence traffic, interarea path optimization, metric engineering.

質問: **141**

展示品を参照してください。



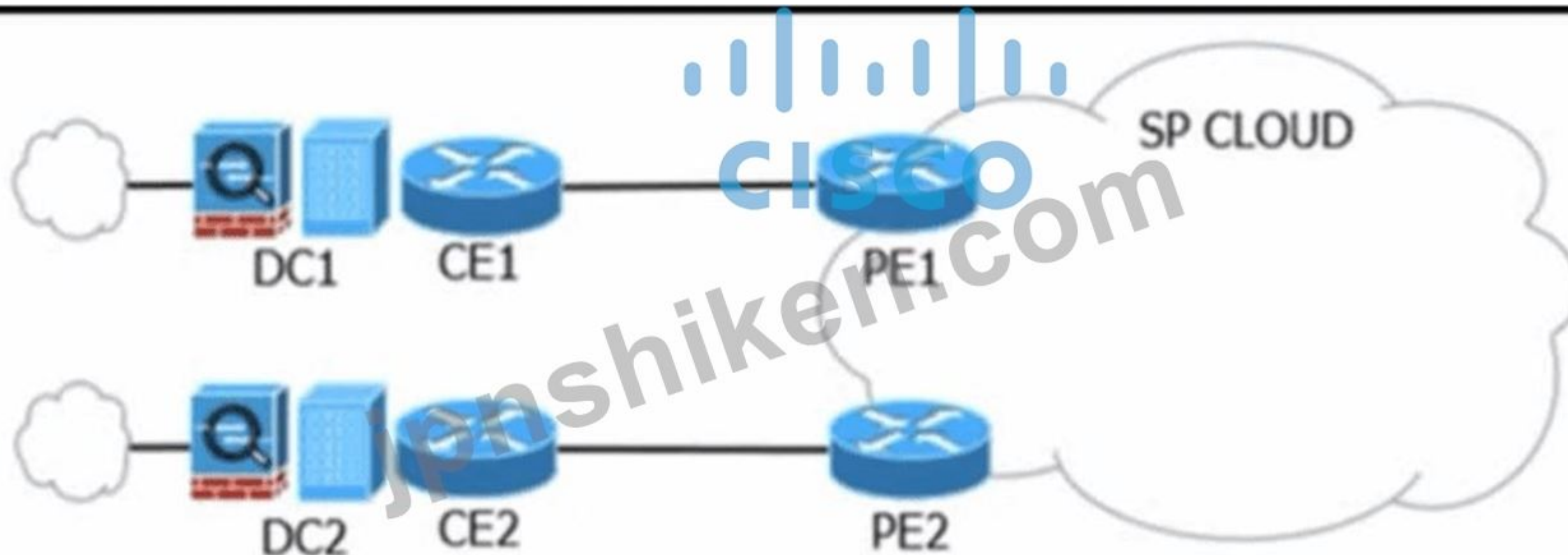
顧客はコア ルータで HSRP を実行しています。時間の経過とともに会社は成長し、より多くのネットワーク容量が必要になりました。現在の環境では、ダウンストリーム インターフェイスの一部はほぼ完全に使用されていますが、他のインターフェイスは使用されていません。どのソリューションが状況を改善しますか。

- A. ルータ R2 を VLAN の半分に対してアクティブにします。
- B. R1 と R2 にさらにインターフェイスを追加します。
- C. ダウンストリーム スイッチへのポート チャネルを設定します。
- D. ダウンストリーム スイッチで RSTP を有効にします。

正解: (正解を表示します)

Making R2 active for half of the VLANs improves utilization by distributing first-hop gateway forwarding across both core routers. HSRP provides default-gateway redundancy by making one router active and another standby for a VLAN. If the same router is active for every VLAN, traffic from all downstream VLANs is forwarded through the same core path, which can leave some interfaces congested while parallel links on the standby side remain underused. Cisco campus designs commonly tune HSRP priorities per VLAN and align the active HSRP gateway with the spanning-tree root for that VLAN. This creates a deterministic load-sharing model while preserving gateway redundancy. Adding more interfaces may help capacity, but the question states that some downstream interfaces are already underutilized, so the issue is load distribution rather than absolute port count. A port channel is useful only when the topology supports bundling the specific links and does not by itself move gateway activity. RSTP improves convergence, not HSRP forwarding balance. Reference topics: HSRP load sharing, per-VLAN gateway placement, STP root alignment, campus capacity design.

質問: 142



図を参照してください。ある企業が顧客向けにアプリケーションを開発しましたが、今度はそれを展開する必要があります。

アプリケーションの展開は、以下の要件を満たす必要があります。

- A. 2つのファイアウォールを接続します。アプリケーションをDC1とDC2にデプロイします。IP SLAを使用してDC2からのアドバタイズメントを制御します。
- B. 2つのファイアウォールを接続します。アプリケーションをDC1とDC2にデプロイします。DC1とDC2から同じプレフィックスをアドバタイズします。
- C. アプリケーションを DC1 と DC2 にデプロイします。DC1 から /32 のプレフィックスをアドバタイズします。DC2 から /24 のプレフィックスをアドバタイズします。

D. アプリケーションを DC1 と DC2 にデプロイします。DC1 と DC2 から同じプレフィックスをアドバタイズします。トラフィック フローを分散します。

正解: **C** ([コメントを发表する](#))

The correct design is to deploy the application in both data centers, advertise the application prefix from DC1 as a /32, and advertise a broader /24 from DC2. This uses longest-prefix-match behavior to create deterministic primary and backup routing. Cisco routing logic prefers the most specific matching route in the routing table, regardless of administrative distance or metric comparisons among less specific covering routes.

As long as the DC1 /32 is present, traffic for that exact application address follows DC1. If DC1, its edge path, or the /32 advertisement fails, the /32 is withdrawn and the remaining /24 covering route from DC2 provides backup reachability. Advertising the same prefix from both locations may result in load sharing or provider-dependent selection, which does not guarantee DC1 primary behavior. IP SLA could be used in some route-tracking designs, but the clean routing solution tested here is prefix specificity. Reference topics:

longest-prefix match, BGP advertisement, active/standby data center routing, covering routes, route failover.

質問: **143**

アーキテクトは、以下を含むキャンパスネットワークソリューションを開発する必要があります。

論理的にセグメント化され分離されたネットワーク

必要に応じてネットワークセグメント間で通信する機能

重複するIPアドレスのサポート

特殊な機器の購入を回避するために広く利用可能なテクノロジーアーキテクトはどのソリューションを選択する必要がありますか？

A. IGPを使用したVSS

B. HSRPを使用した802.1Q

C. HSRPを使用したvPC

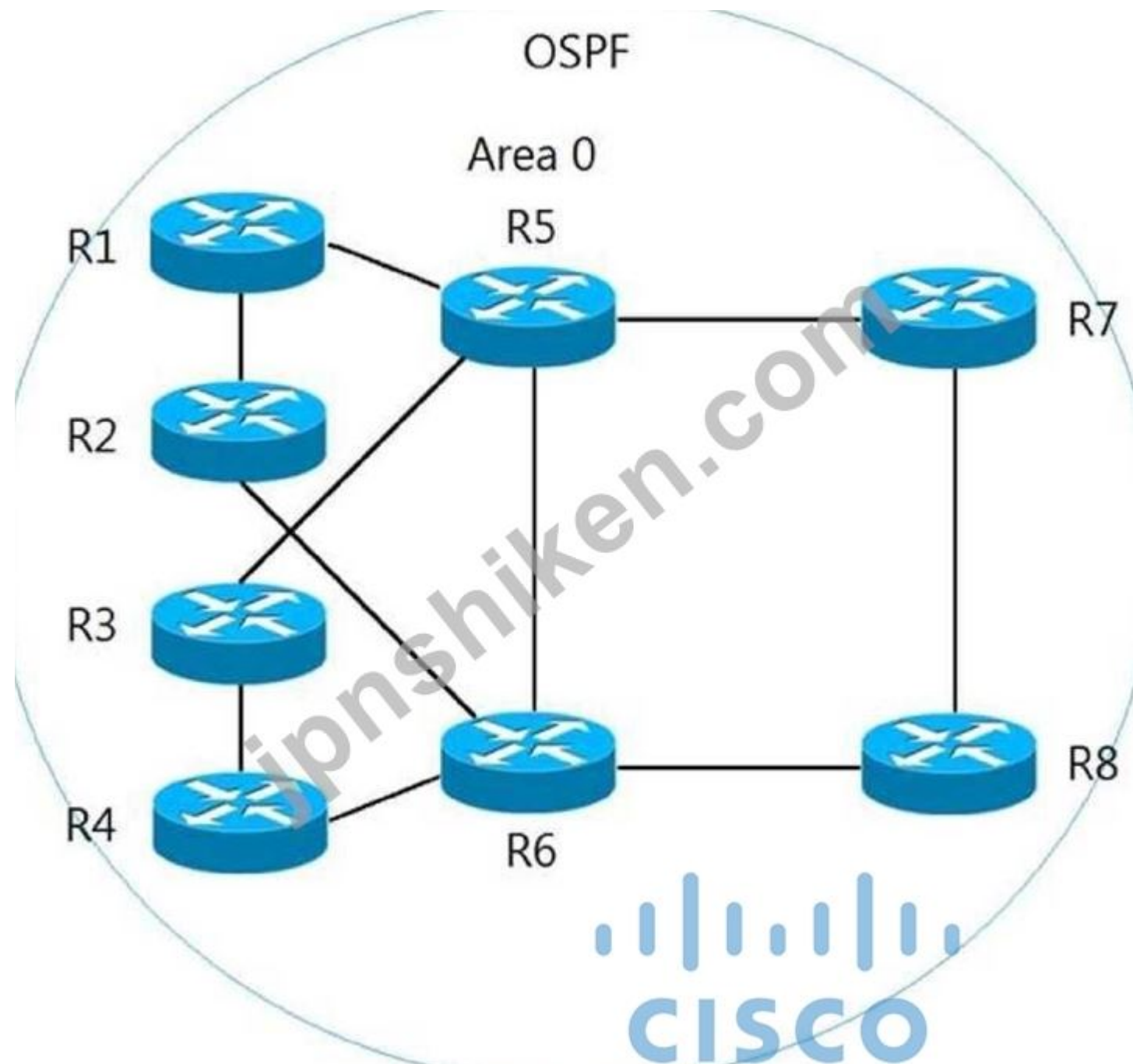
D. OSPFを使用したVRF-Lite

正解: ([正解を表示します](#))

VRF-Lite with OSPF satisfies the requirements because it creates multiple isolated Layer 3 routing tables without requiring a service-provider MPLS core or specialized hardware. Each VRF has its own forwarding table, interfaces, and route exchange process, so overlapping IP address space can exist in different virtual networks. When communication between segments is required, the design can use controlled route leaking, a firewall, or a shared-services VRF. OSPF can run per VRF to exchange routes within each isolated segment using widely supported routing technology. VSS, vPC, and plain 802.1Q with HSRP can improve Layer 2 or device redundancy, but they do not provide independent Layer 3 routing tables with overlapping addressing.

A multihop MPLS design can also support many VPNs, but it is more complex and beyond the stated need for widely available technologies without specialized equipment. VRF-Lite is therefore the correct enterprise campus segmentation choice. Reference topics: VRF-Lite, virtual routing and forwarding, inter-VRF routing, overlapping IP address support, campus segmentation.

質問: **144**



図を参照してください。現在、すべてのルータは OSPF エリア 0 に存在します。ネットワーク マネージャは最近、R1 と R2 をリモート ブランチ ロケーションの集約ルータとして使用し、R3 と R4 をリモート オフィス ロケーションの集約ルータとして使用しました。それ以来、ネットワークは停止に悩まされており、SPF の実行が頻繁に発生しています。安定性を高め、可能な限り最小限の ABR で OSPF ネットワークにエリアを導入するには、ネットワーク マネージャが推奨する 2 つのソリューションはどれですか。(2 つ選択してください。)

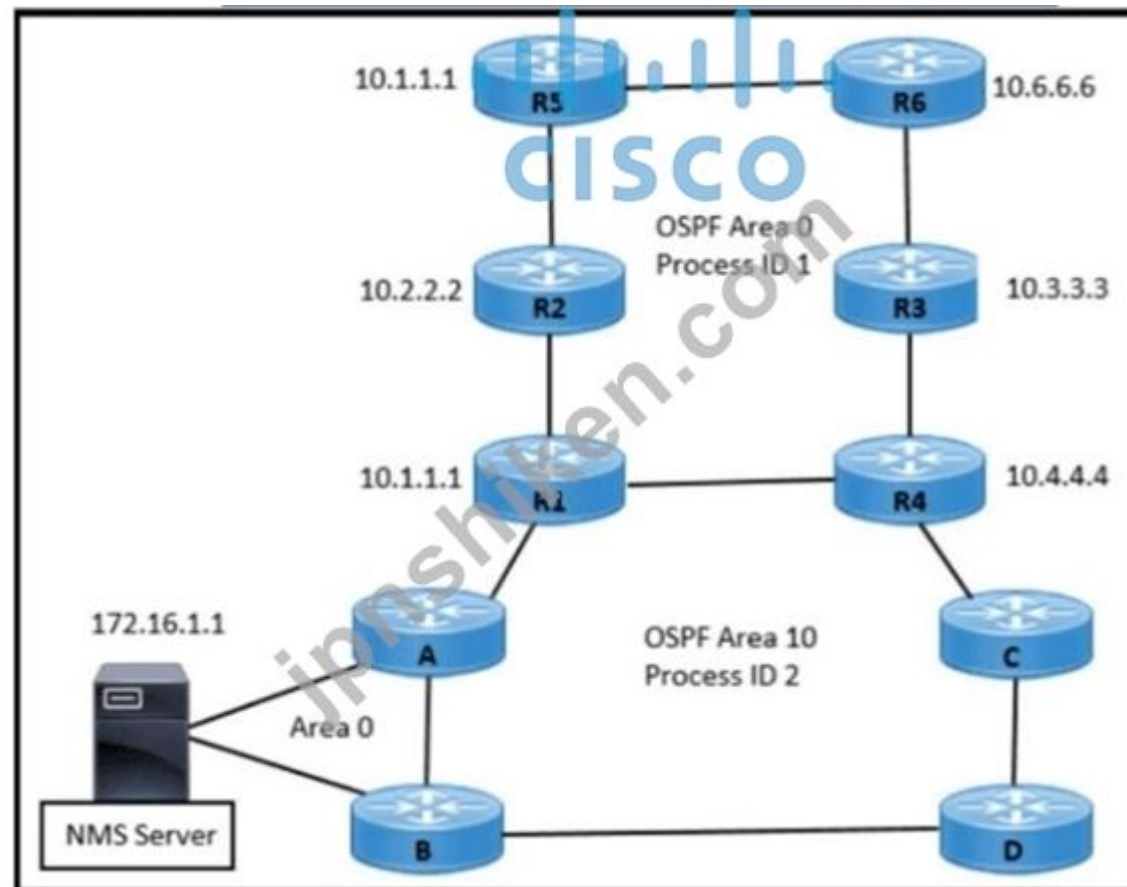
- A. R1とR2をABRとするR1とR2接続用の新しいOSPFエリア
- B. R3 と R4 接続用の新しい OSPF エリア。R5 と R6 は ABR として使用されます。
- C. R3とR4をABRとするR3とR4接続用の新しいOSPFエリア
- D. R1、R2、R3、R4 接続用の新しい OSPF エリア。R1、R2、R3、R4 は ABR として使用されます。
- E. R1 と R2 接続用の新しい OSPF エリア。R5 と R6 は ABR として使用されます。

正解: **B,E** ([コメントを发表する](#))

The design goal is to reduce the impact of branch and remote-office instability while using the fewest possible ABRs. In OSPF, topology changes inside an area trigger SPF calculations for routers in that area. Creating separate areas behind aggregation points contains those changes and prevents every branch link flap from affecting the entire backbone. The minimal-ABR approach is to place the area

boundary at the higher aggregation layer, where fewer routers connect the new areas to area 0. In the described topology, R5 and R6 can serve as ABRs for both sets of downstream aggregation connections. That supports a new area for the R1/R2 branch connections with R5 and R6 as ABRs, and a new area for the R3/R4 remote-office connections with R5 and R6 as ABRs. Making R1, R2, R3, and R4 the ABRs increases the ABR count and spreads area- boundary functions deeper into the network. A single new area for all downstream connections would reduce isolation between different failure domains. The selected design improves stability and keeps the OSPF hierarchy clean.

質問: 145



展示を参照してください。エンジニアは、次の要件を満たす OSPF ソリューションを設計しています。

* NMS サーバーは R5 と R6 を管理します。

* R1 に障害が発生した場合、すべての NMS トラフィックは R4 を経由してルーティングされる必要があります。

* R5 と R6 間のリンクに障害が発生した場合、10.6.6.6 宛てのすべてのトラフィックは R4 を経由してルーティングされる必要があります。エンジニアはどのソリューションを選択する必要がありますか？

A. R1 上の高コストで 172.16.1.1 を OSPF プロセス 1 にアドバタイズします。

B. R5 および R6 への IP SLA トラッキングを使用して、R2 および R3 に静的ルートを適用します。

C. R2 から R1 へのより高いメトリックを使用して、default-information originate コマンドを有効にします。

D. OSPF プロセス 1 を R1 および R4 上のプロセス 2 に再配布します。

正解: (正解を表示します)

Redistributing OSPF process 1 into process 2 on both R1 and R4 provides the required backup routing behavior between the management network and the target routers. The design requirements are failure driven:

if R1 fails, NMS traffic must use R4, and if the link between R5 and R6 fails, traffic to 10.6.6.6 must still route through R4. When separate OSPF processes or domains exist, controlled redistribution at both exit points allows reachability to be maintained through either redistribution point. Advertising a single address into OSPF with high cost on R1 does not solve both failure cases. Static routes with IP SLA on R2 and R3 would be more manual and less aligned with the dynamic OSPF design. Default-information originate would provide a default path but would not accurately advertise the specific internal

reachability required for the NMS destinations. The redistribution design must also include route filtering or tagging in production to prevent feedback loops. Reference topics: OSPF redistribution, redundant redistribution points, management reachability, route control, failure-path design.

質問: 146

ネットワークでモデル駆動型テレメトリを使用する利点は何ですか？

- A. 割り込み駆動型ポーリングを使用して、一定の間隔でデータを取得します。
- B. JSON エンコーディングを使用し、市場のさまざまなツールと互換性があります。
- C. 業界でよく知られているデータを構造化するために MIB モデルを使用します。
- D. テレメトリは、show コマンドからの CLI 出力を解析してデータを取得します。

正解: ([正解を表示します](#))

A practical advantage of model-driven telemetry is that it sends structured data in machine-readable encodings, including JSON in supported implementations, making the data easier to consume with modern collectors and automation tools. Model-driven telemetry uses data models, commonly YANG, to define exactly which operational or configuration data is streamed. This is a major improvement over legacy monitoring approaches that repeatedly poll devices or parse human-oriented CLI output. The question option that points to structured JSON compatibility is the best available answer. Telemetry is not based on interrupt- driven polling; it is normally subscription based, either periodic or on-change. It also does not rely on MIB models as its primary structure in the way SNMP does. Parsing show command output is the older, brittle method that model-driven telemetry is intended to replace. The design benefit is cleaner, more predictable, and more programmable network state collection. Reference topics: model-driven telemetry, YANG data models, JSON encoding, subscription-based monitoring, network programmability.

質問: 147

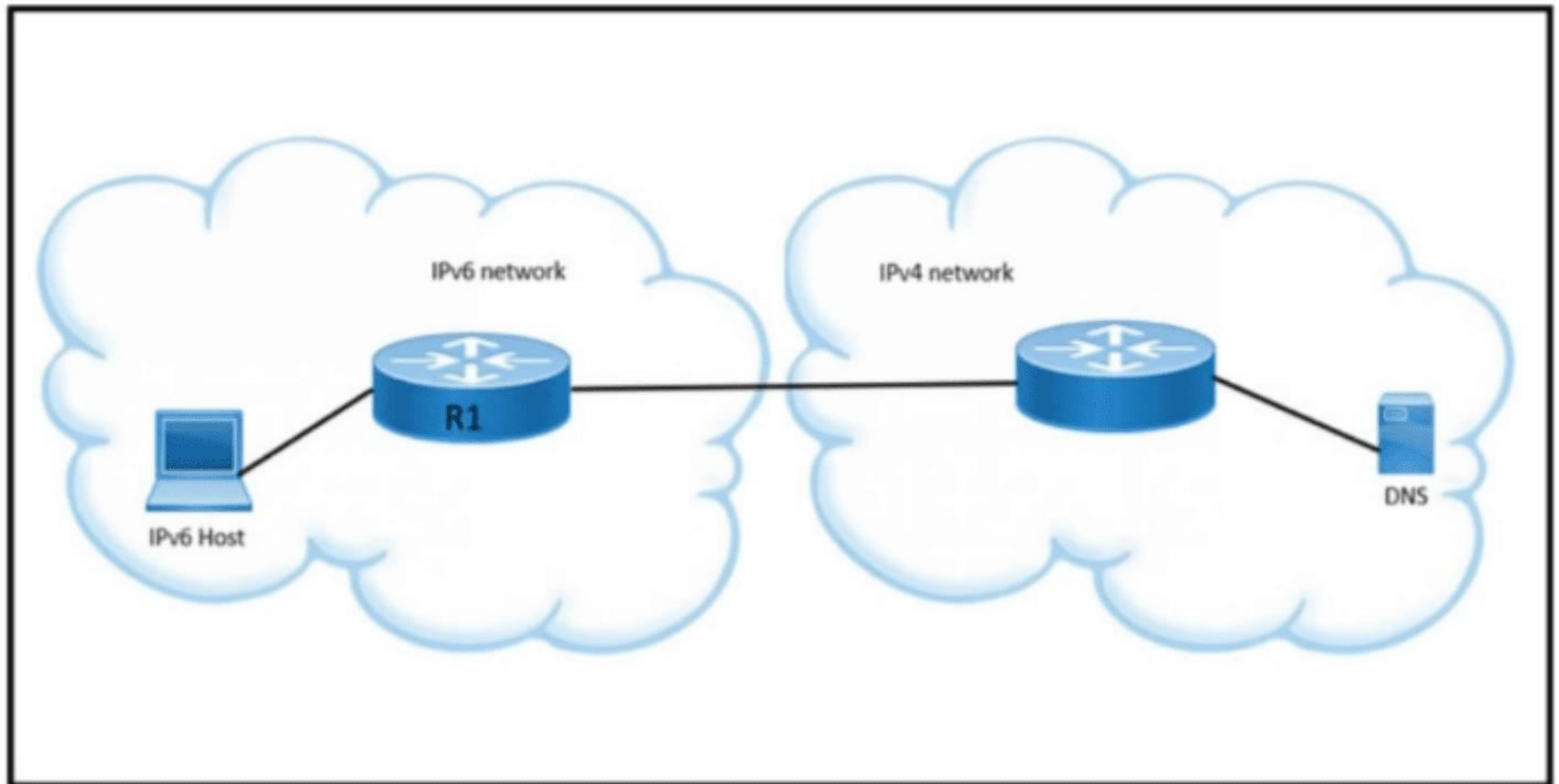
モデル駆動型のテレメトリ ダイヤルアウト アプローチはどのように機能しますか？

- A. デバイスは、サブスクリプションに基づいてコレクターへのセッションを開始します。
- B. コレクターはデバイスへのセッションを開始し、ストリーミングされるデータをサブスクライブします。
- C. コレクターは、デバイスへのセッションを開始し、以前に定義されたサブスクリプションのデータを取得します。
- D. デバイスはコレクターへのセッションを開始し、サブスクリプションをネゴシエートします。

正解: **A** ([コメントを發表する](#))

The dial-out model-driven telemetry approach is initiated by the network device, not by the collector. Cisco IOS XE telemetry documentation describes configured subscriptions as dial-out because the publisher, meaning the router or switch, initiates the connection to the configured receiver. After the session is established, the device streams the subscribed YANG-modeled data at the configured cadence or when the configured on-change condition occurs. This model is useful when the telemetry collector is a fixed destination and the network operator wants the device to maintain the outbound session automatically. Option B describes dial-in telemetry, where the collector or subscriber connects to the device and requests the subscription. Option C is also dial-in oriented because the collector initiates the session. Option D is close in wording, but the key design point is that the session is based on an already configured subscription rather than a negotiation process. Reference topics: model-driven telemetry, configured subscriptions, dial-out telemetry, YANG-push, telemetry collectors.

質問: 148



図を参照してください。エンジニアは、IPv6 アイランドを IPv4 のみのネットワークに接続して、IPv6 ホストが IPv4 ネットワーク内のファイル サーバーと DNS サービスにアクセスできるようにする必要があります。エンジニアはどの NAT を選択する必要がありますか？

- A. ステートレス NAT66
- B. ステートフル NAT66
- C. 静的NAT-PT
- D. 動的NAT-PT

正解: [\(正解を表示します\)](#)

Dynamic NAT-PT is the correct selection when IPv6 hosts in an IPv6 island must access IPv4-only servers and services. NAT-PT translates between IPv6 and IPv4 packet headers and allows IPv6-only nodes to communicate with IPv4-only destinations by dynamically creating translation state. In this scenario, IPv6 hosts need access to file servers and DNS services in an IPv4 network, so a dynamic translation mechanism is required. Stateless NAT66 and stateful NAT66 are IPv6-to-IPv6 translation mechanisms and do not connect IPv6 hosts to IPv4-only resources. Static NAT-PT would be appropriate

for fixed one-to-one mappings where specific hosts are permanently mapped, but the question describes general access from an IPv6 island to IPv4 services, making dynamic NAT-PT the better match. In modern designs, NAT64 and DNS64 are commonly preferred successors, but within the answer choices NAT-PT is the IPv6-to-IPv4 transition mechanism.

Reference topics: IPv6 transition, NAT-PT, IPv6-to-IPv4 translation, DNS access, dual-protocol migration.

質問: 149

ネットワークエンジニアが、Ciscoレイヤ2スイッチの代替ポートまたはルートポートが断続的に指定ポートになり、STPループが発生することを発見しました。この問題を解決するには、どのような設定が必要ですか？

A. PortFast BPDUガード

B. UDLD

C. STPループガード

D. STPルートガード

正解: ([正解を表示します](#))

STP loop guard is the correct protection feature. Cisco defines loop guard as a spanning-tree enhancement that prevents alternate or root ports from incorrectly becoming designated forwarding ports when expected BPDUs are no longer received. That failure mode is exactly what the question describes: an alternate or root port intermittently transitions to a designated role and creates a Layer 2 loop. With loop guard enabled, the port is placed into a loop-inconsistent blocking state until BPDUs are received again, so the topology remains protected. PortFast BPDU guard is intended for edge ports connected to end hosts and shuts a PortFast port when BPDUs appear; it is not the correct uplink/root-port protection. UDLD can detect unidirectional links, but it does not specifically enforce STP role consistency. Root guard prevents superior BPDUs from changing root placement, which is a different problem. Reference topics: STP loop guard, alternate and root ports, BPDU loss, loop-inconsistent state, Layer 2 loop prevention.

質問: 150

あるアーキテクトが、大規模なデータ センターのインフラストラクチャを会社がどのように管理するかを設計しています。この会社は、セキュリティ上の理由からデバイスの種類をグループ化し、DoS 攻撃を軽減したいと考えています。また、管理プレーンで 1 つのデバイスが侵害された場合、運用ネットワークの残りの部分にアクセスできないようにしたいと考えています。アーキテクトはどのソリューションを選択する必要がありますか？

A. インバンドダイヤルアップ回線

B. インバンドイーサネット

C. 帯域外イーサネット

D. 帯域外ダイヤルアップ回線

正解: ([正解を表示します](#))

Out-of-band Ethernet is the correct management design for a large data center when the customer wants device grouping, management-plane isolation, and reduced blast radius from compromise. In an out-of-band design, management traffic uses a physically or logically separate management network instead of traversing the production data plane. This allows different device classes to be placed into management zones, protected by access control, authentication, logging, and rate-limiting policies. If one managed device is compromised, the attacker should not automatically gain production network reachability through the management interface.

In-band Ethernet would share the production infrastructure and therefore would not provide the same isolation. Dial-up out-of-band access can be useful for emergency console access, but it is not appropriate as the primary management architecture for a large data center. Out-of-band Ethernet provides scale, speed, and segmentation while still allowing centralized management systems to reach routers, switches, controllers, and appliances. Reference topics: out-of-band management, management-plane security, data-center operations, isolation, access control, DoS mitigation.

質問: 151

ブランチ オフィスには、メイン オフィスに戻るプライマリ L3VPN MPLS 接続と、バックアップとして機能する IPSEC VPN トンネルがあります。プライマリ MPLS 回線がダウンした場合にのみ、データがバックアップ接続を介して送信されるようにする設計はどれですか。

A. EIGRPを使用して、本社とのネイバー関係を確立します。

B. L3VPN MPLS と IPSEC VPN トンネル。

C. マルチパス機能を有効にした BGP を使用して、使用可能な場合はトラフィックをプライマリ パス経由で強制します。

D. フローティング スタティック ルートがバックアップ接続を指しているときに、IP SLA に関連付けられたスタティック ルートを使用してプライマリ パスを優先します。

E. バックアップ接続で、passive-interface コマンドを使用して OSPF を使用します。

正解: D ([コメントを发表する](#))

Static routes tied to IP SLA with a floating static backup route provide the cleanest active/standby design for this requirement. The MPLS L3VPN path is the primary connection and should carry traffic while it is available. IP SLA can actively test reachability across the primary path, and object tracking can remove the primary static route if the test fails. The backup IPsec VPN route is configured with a higher administrative distance, so it remains inactive until the primary route is withdrawn. This design is deterministic and prevents traffic from using the backup tunnel during normal operation. Running EIGRP across both paths could allow dynamic path selection, but it would require careful metric manipulation and could still introduce unintended failover behavior. BGP multipath is the opposite of the requirement because it is designed to install multiple paths simultaneously for load sharing. OSPF passive-interface on the backup connection would prevent neighbor formation rather than provide reliable backup routing. For a branch with one primary MPLS path and one backup IPsec path, tracked static routing with IP SLA is the appropriate design.

有効的な**300-420J**問題集はJPNTTest.com提供され、**300-420J**試験に合格することに役に立ちます！JPNTTest.comは今最新**300-420J**試験問題集を提供します。JPNTTest.com 300-420J試験問題集はもう更新されました。ここで**300-420J**問題集のテストエンジンを手に入れます。最新版のアクセス、<https://www.jpntest.com/shiken/300-420J-mondaishu> **457**問、**30%ディスカウント**、特別な割引コード：**JPNshiken**」

質問: 152

エンジニアは、XML 表現で YANG を使用して、次の仕様で Cisco IOS XE スイッチを設定する必要があります。

直接接続されたホスト 10.10.10.1/27 からのインターフェイス GigabitEthernet2/1/0 接続で設定された IP アドレス 10.10.10.10/27 エンジニアが選択する必要がある YANG データ モデル セットはどれですか。

```
<interfaces xmlns="urn:ietf:params:xml:ns:yang:ietf-interfaces">
  <interface>
    <name>GigabitEthernet2/1/0</name>
    <type xmlns:ianaift="urn:ietf:params:xml:ns:yang:iana-if-type">ianaift:ethenetCsmacd</type>
    <enabled>>false</enabled>
    <ipv4 xmlns="urn:ietf:params:xml:ns:yang:ietf-ip">
      <address>
        <ip>10.10.10.10</ip>
        <netmask>255.255.255.224</netmask>
      </address>
    </ipv4>
  </interface>
</interfaces>
```

A.

```
<interfaces YANG="urn:ietf:params:xml:ns:yang:ietf-interfaces">
  <interface>
    <name>GigabitEthernet2/1/0</name>
    <type YANG:ianaift="urn:ietf:params:xml:ns:yang:iana-if-type">ianaift:ethernetCsmacd</type>
    <enabled>true</enabled>
    <ipv4 YANG="urn:ietf:params:xml:ns:yang:ietf-ip">
      <address>
        <ip>10.10.10.10</ip>
        <netmask>255.255.255.224</netmask>
      </address>
    </ipv4>
  </interface>
</interfaces>
```

B.

```
<interfaces json="urn:ietf:params:json:ns:yang:ietf-interfaces">
  <interface>
    <name>GigabitEthermet2/1/0</name>
    <type json:ianaift="urn:ietf:params:json:ns:yang:iana-if-type">ianaift:ethernetCsmacd</type>
    <enabled>true</enabled>
    <ipv4 json="urn:ietf:params:json:ns:yang:ietf-ip">
      <address>
        <ip>10.10.10.10</ip>
        <netmask>255.255.255.224</netmask>
      </address>
    </ipv4>
  </interface>
</interfaces>
```

C.

```

<interfaces xmlns="urn:ietf:params:xml:ns:yang:ietf-interfaces">
  <interface>
    <name>GigabitEthernet2/1/0</name>
    <type xmlns:ianaift="urn:ietf:params:xml:ns:yang:iana-if-type">ianaift:ethernetCsmacd</type>
    <enabled>true</enabled>
    <ipv4 xmlns="urn:ietf:params:xml:ns:yang:ietf-ip">
      <address>
        <ip>10.10.10.10</ip>
        <netmask>255.255.255.224</netmask>
      </address>
    </ipv4>
  </interface>
</interfaces>

```

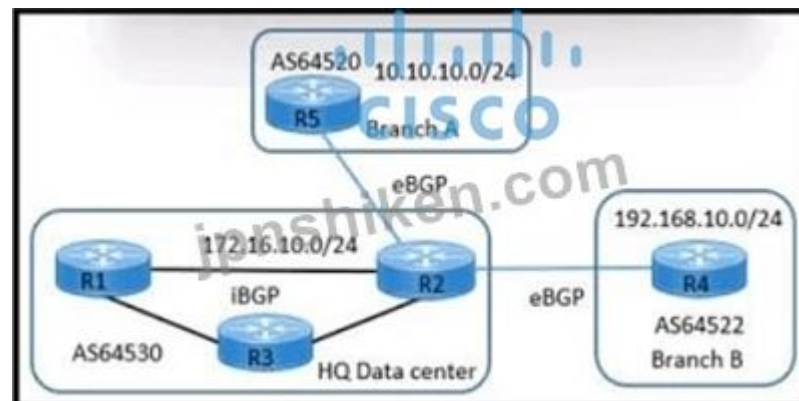
D.

正解: ([正解を表示します](#))

The correct XML/YANG data model set is the option that accurately represents the IOS XE interface hierarchy, assigns the IPv4 address to GigabitEthernet2/1/0, and uses the correct prefix length of 27. In model-driven programmability, the payload must match both the selected YANG model schema and the target platform's interface naming conventions. For IOS XE native models, interface configuration is structured under the Cisco native namespace, with interface type and name represented in the schema before IPv4 address elements are applied. The address 10.10.10.10/27 must also support a directly connected host at

10.10.10.1/27, which places both addresses in the same 10.10.10.0/27 subnet. Options that place the address under the wrong interface type, use the wrong namespace, incorrectly represent the mask, or configure the host route rather than the interface are invalid. Option D is the only answer aligned to the IOS XE YANG structure and the requested addressing. Reference topics: IOS XE native YANG, NETCONF XML payloads, interface configuration, IPv4 prefix length, model-driven programmability.

質問: 153



展示を参照してください。従業員 ID: 4384:99:754 のネットワーク エンジニアは、次の条件に基づいて BGP ソリューションを設計する必要があります。

- * トラフィック セッションはブランチとデータ センター間で発生します。
- * ブランチ B にはルーティング更新を処理するためのリソースが限られています。
- * HQ はブランチ A から R4 までのすべてのプレフィックスをフィルタリングする必要があります。

エンジニアはどのアウトバウンド ルート フィルタリング (ORF) ソリューションを選択する必要がありますか？

- A. R4 上の ORF に 192.168.10.0/24 サブネットのプレフィックス リストを使用します。
- B. R2のORFに10.10.10.0/24サブネットのプレフィックスリストを使用する
- C. R5 の ORF に 10.10.10.0/24 サブネットのプレフィックス リストを使用します。
- D. R2 上の ORF に 192.168.10.0/24 サブネットのプレフィックス リストを使用します。

正解: ([正解を表示します](#))

The ORF filter should be applied so the route reflector or upstream peer suppresses the unwanted Branch A prefix before sending updates toward the resource-constrained branch. Cisco BGP outbound route filtering lets a receiving router advertise a prefix-list filter to its peer, causing that peer to apply the filter outbound.

This is useful when a router has limited CPU or memory because unwanted prefixes are never sent across the session. In the described design, HQ must prevent Branch A prefixes from reaching R4, while Branch B has limited resources. The answer using the 10.10.10.0/24 prefix list on R2 aligns with filtering the Branch A prefix at the point where HQ can control route reflection or advertisement toward the constrained branch path.

Applying the filter on the wrong router or using the wrong prefix would still allow unnecessary updates. ORF is preferred over simple inbound filtering when the goal is to reduce received route volume.

Reference topics:

BGP ORF, prefix-list signaling, route-reflector filtering, branch route scale, outbound update suppression.

質問: **154**

ネットワークエンジニアは、マルチキャストトラフィックをサポートする安全なトンネリングテクノロジーを使用して、パブリックネットワークを介して2つのサイトを接続する必要があります。どのテクノロジーを選択する必要がありますか？

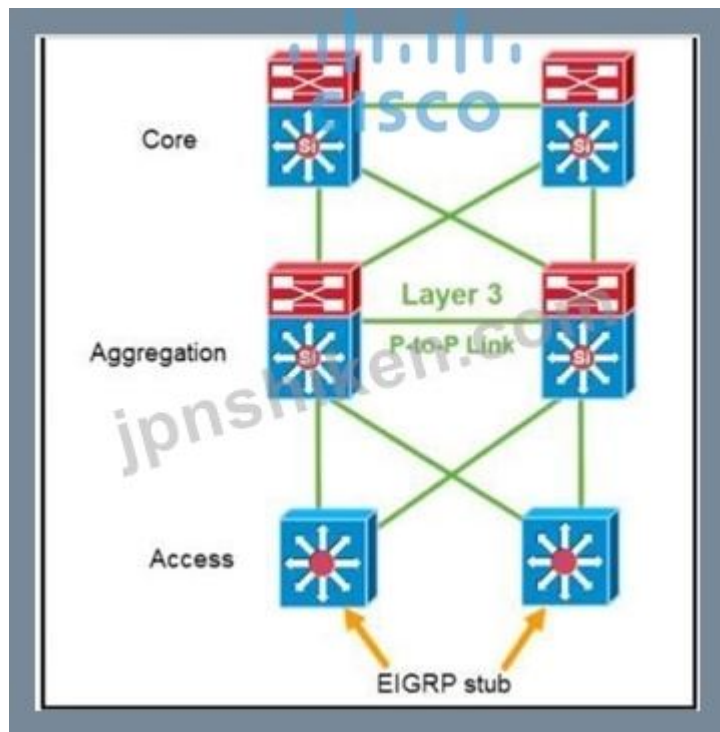
- A. IPsec
- B. GRE
- C. PPTP
- D. GRE over IPsec

正解: ([正解を表示します](#))

GRE over IPsec is the correct secure tunneling choice when multicast traffic must cross a public network.

Native IPsec protects unicast IP traffic, but it does not inherently carry multicast routing protocols or multicast application flows in the same way that GRE can. GRE creates a logical point-to-point tunnel that can encapsulate multicast, broadcast, and non-IP payloads; IPsec then encrypts the GRE packet for confidentiality and integrity across the untrusted transport. Cisco designs commonly use this combination when dynamic routing protocols, multicast applications, or other non-unicast traffic must be transported securely between sites. GRE by itself supports multicast, but it does not provide encryption, so it fails the secure tunneling requirement. PPTP is obsolete and not appropriate for modern enterprise design. Plain IPsec is secure, but it is not the complete answer when multicast support is explicitly required. Therefore, the design should use GRE over IPsec: GRE supplies the multicast-capable overlay, and IPsec supplies the cryptographic protection across the public network.

質問: **155**



図を参照してください。アーキテクトはトポロジのルート集約をどこで計画する必要がありますか？

- A. コアから集約へ、そして集約へのアクセスへ
- B. コアから集合体へ、集合体からコアへ
- C. 集約からアクセスへ、そしてアクセスから集約へ
- D. 集約からコアへ、集約からアクセスへ

正解: ([正解を表示します](#))

Route summarization should be planned from the aggregation layer toward both the core and the access layer.

In hierarchical campus routing, aggregation or distribution devices are natural summarization boundaries because they sit between more specific access-layer networks and the core. Summarizing from aggregation toward the core prevents access-layer subnet churn from forcing unnecessary recomputation and route table growth in the core. Summarizing from aggregation toward access can also provide a concise upstream route, often a default or summary route, so access devices do not need the entire network routing table. The goal is to contain failures, reduce routing state, and preserve fast convergence at hierarchy boundaries. Summarizing from the core toward aggregation may hide useful topology details in the wrong direction, while summarizing directly from access toward aggregation is not always possible or efficient when access devices are numerous and have limited resources. The answer that places summarization at aggregation in both directions best matches Cisco hierarchical design principles. Reference topics: hierarchical campus design, route summarization, aggregation layer, core route scale, convergence containment.

質問: 156

ネットワークエンジニアは、IS-ISルーティングを使用して相互接続された3つのキャンパスネットワークを分離する必要があります。大規模なルーティングドメインをサポートし、各キャンパスネットワークから他のキャンパスネットワークルーターに自動的にアドバタイズされるより具体的なルート回避するには、2層の階層を使用する必要があります。この分離を達成するためにエンジニアはどの2つのアクションを実行しますか？ 2つ選択してください。）

- A. 各キャンパスのエッジで2つのIS-ISルーターをBDRルーターとして指定し、すべてのレベル1ルーターに1つのBDRを構成し、すべてのレベル2ルーターに1つのBDRを構成します。
- B. 各キャンパスネットワークのエッジでレベル1/レベル2バックボーンルーターとして機能するように、各キャンパスから2つのIS-ISルーターを指定します。
- C. 各キャンパスに同じIS-IS NET値を割り当て、レベル1/レベル2ルーティングで内部キャンパスルーターを構成します。
- D. キャンパスネットワークセグメントごとに異なるMTU値を利用します。レベル2バックボーンルーターは、9216のより大きなMTUサイズを利用する必要があります。
- E. キャンパスごとに一意のIS-IS NET値を割り当て、レベル1ルーティングで内部キャンパスルーターを構成します。

正解: ([正解を表示します](#))

A scalable two-level IS-IS design uses Level 1 areas for internal campus routing and Level 2 for the interarea backbone. Each campus should have a unique IS-IS NET area value so the internal routers belong to that campus's Level 1 area. Internal campus routers should be configured as Level 1 routers, which keeps detailed campus prefixes inside the local area and prevents unnecessary propagation of more-specific routes to other campuses. At the edge of each campus, routers that connect the campus to the backbone should operate as Level 1/Level 2 routers. These routers connect the local Level 1 area to the Level 2 backbone and can summarize or control routing information between levels. IS-IS does not use BDR routers; that is an OSPF term, so option A is invalid. Assigning the same NET value everywhere collapses the hierarchy and defeats the segregation goal. MTU differences are not a hierarchy or route-segmentation tool. Therefore, the design should use unique NET values per campus with Level 1 internal routers and Level 1/Level 2 routers at the campus edge.

質問: 157

エンジニアは、BGP を使用してインターネット接続のために ISP とピアリングする必要があります。最初は、エンジニアは ISP から特定のプレフィックスとデフォルト ルートのみを受け取りたいと考えています。ただし、このソリューションは、将来的にすぐにプレフィックスを追加できる柔軟性を提供する必要があります。ISP では、2 週間の変更プロセスが用意されています。エンジニアはどのルートフィルタリング ソリューションを採用する必要がありますか？

- A.** 制限付きのインターネット ルーティング テーブルとデフォルト ルートを ISP に要求し、特定のインターネット プレフィックスとブロックされたネットワークのみを許可するアクセス リストを使用して、BGP の最大制限を 1 に設定します。
- B.** ホワイトリストに登録されたネットワークを使用して、必要なプレフィックスとデフォルト ルートのみを ISO からアドバタイズするように要求します。
- C.** ISP から完全なインターネット ルーティング テーブルとデフォルト ルートを要求し、デフォルト ルートと必要なプレフィックスを許可するプレフィックス リストを使用してインバウンド ルート フィルタリングを構成します。
- D.** 必要なプレフィックスを企業が ISP に伝えるように、企業と ISP でアウトバウンド ルート フィルタリングを構成します。

正解: [\(正解を表示します\)](#)

BGP outbound route filtering is the correct solution because the enterprise wants control over which routes it receives without waiting for a provider change cycle each time the prefix set changes. Cisco BGP prefix- based ORF allows a receiving BGP speaker to send an inbound prefix-list filter to the remote peer, which then applies that list as an outbound filter. This reduces unwanted routing updates before they cross the session and lets the enterprise update the desired prefixes locally when business requirements change. Requesting only fixed prefixes from the ISP would work initially but fails the flexibility requirement because the provider has a two-week change process. Requesting a full table and filtering inbound would provide flexibility but wastes memory and CPU on the customer router because all routes are still received before filtering. A maximum- prefix limit does not select the correct routes. ORF is therefore the scalable operational design for receiving a default route and selected prefixes with minimal provider involvement. Reference topics: BGP ORF, prefix- based filtering, inbound route control, ISP peering design, routing table scale.

質問: 158

PIM スパース モードで、マルチキャスト対応デバイスで RPF チェックが成功した場合、マルチキャスト パケットはどうなりますか？

- A.** OIL 内のすべてのインターフェイスに転送されます。
- B.** 受信インターフェースを除くすべてのインターフェイスに転送されます。
- C.** 転送されたパケットはループを防ぐためにドロップされます。
- D.** PIM 対応のすべてのインターフェイスに転送されます。

正解: **A** ([コメントを發表する](#))

If the RPF check is successful in PIM sparse mode, the multicast packet is forwarded according to the outgoing interface list. The Reverse Path Forwarding check verifies that the packet arrived on the interface the router would use to reach the source, or the RP in the case of shared-tree state. This prevents loops and duplicate multicast forwarding. After the check passes, the router consults the multicast forwarding entry and sends the packet only out interfaces in the OIL where downstream receivers or PIM neighbors require the traffic. It is not automatically forwarded out every PIM-enabled interface, and it is not forwarded back out the receiving interface. If the RPF check fails, the packet is dropped to protect the multicast distribution tree. The OIL therefore represents the correct controlled forwarding behavior after a successful RPF validation.

Reference topics: PIM sparse mode, RPF check, outgoing interface list, multicast forwarding state, loop prevention.

質問: 159

建築家は、地域医療センター向けに、地域に分散したデータセンターと新しいコロケーション施設との相互接続を提供するネットワークソリューションを設計する必要がある。設計には以下の要件が求められる。

- * ポイントツーポイント接続を利用する
- * 既存のVLANインフラストラクチャを活用する
- * データセンターの同期およびバックアッププロセスのパフォーマンスを向上させる
- * 設定の複雑さを軽減する

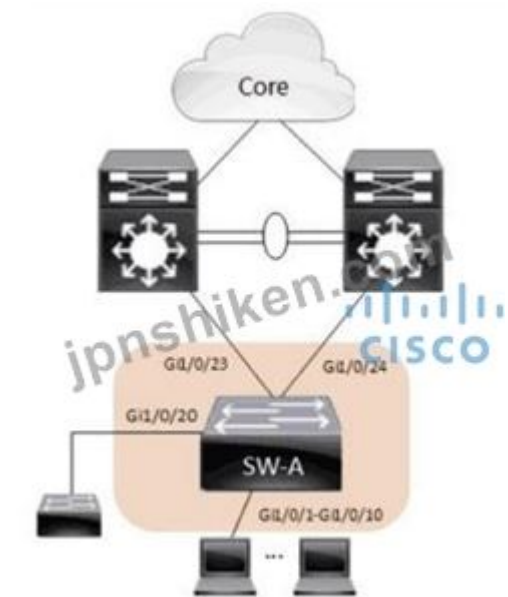
エンジニアはどちらの解決策を選択すべきか？

- A. L3VPN
- B. GO
- C. DMVPN
- D. L2VPN

正解: **D** ([コメントを发表する](#))

L2VPN is the correct solution because the customer requires point-to-point connectivity while preserving existing VLAN infrastructure. Cisco describes Layer 2 VPN services as emulating Layer 2 connectivity across a provider IP or MPLS backbone, commonly through Ethernet pseudowires or VPWS. This allows sites such as dispersed data centers and a colocation facility to exchange Ethernet frames as if they were directly connected at Layer 2. The design is especially useful when data center synchronization, backup workflows, or application dependencies require VLAN continuity or minimal changes to existing addressing and routing. L3VPN would provide routed connectivity and VRF separation, but it would not preserve the VLAN-based service model. GRE would add an overlay tunnel but would not natively provide provider- managed Layer 2 service. DMVPN is designed for scalable Layer 3 IPsec/GRE overlays between branches, not for VLAN extension or point-to-point data center interconnect. Reference topics: L2VPN, VPWS, Ethernet pseudowire, VLAN extension, data center interconnect, WAN service design.

質問: **160**



展示を参照してください。アーキテクトは、会社のエンタープライズ ネットワークの低レベル設計を確認し、STP 収束時間を最適化するようにアドバイスします。アーキテクトの推奨事項に従うには、どの機能を Gi1/0/1-10 に追加する必要がありますか。

- A. ポートファスト
- B. ルートガード
- C. アップリンクファースト
- D. BPDUガード

正解: ([正解を表示します](#))

PortFast is the correct functionality for the Gi1/0/1-10 interfaces when those ports are endpoint-facing access ports. Cisco spanning-tree design uses PortFast to allow edge ports to transition directly to the forwarding state, avoiding the normal listening and learning delay associated with traditional STP convergence. That improves user and endpoint connectivity after link-up events and is especially important for workstations, phones, printers, and access devices that should not participate in the Layer 2 topology. Root guard and BPDU guard are protection features, not the primary convergence optimization requested by the architect.

UplinkFast was used in older STP designs to accelerate recovery for access switches with redundant uplinks, but it is not the edge-port functionality for interfaces Gi1/0/1-10. In a modern design, PortFast should be enabled only on true host-facing ports, normally together with BPDU guard to protect against accidental switch attachment. If the exhibit shows these interfaces as end-user access ports, PortFast is the feature that follows the recommendation to optimize STP convergence time. Reference topics: STP PortFast, edge-port convergence, BPDU guard, Layer 2 campus design.

質問: 161

企業は、建物の 1 つのフロアでアクセス ポートの容量を増やす必要があります。彼らは、既存の触媒アクセス スイッチを活用したいと考えています。アップリンクの帯域容量に問題はありません。ただし、ディストリビューション スイッチで利用できるポートがないため、追加のアップリンクを追加することはできません。追加のアクセス ポートを提供するには、どのソリューションを選択する必要がありますか？

- A. VDC
- B. VSS
- C. イーサチャネル
- D. スタックワイズ

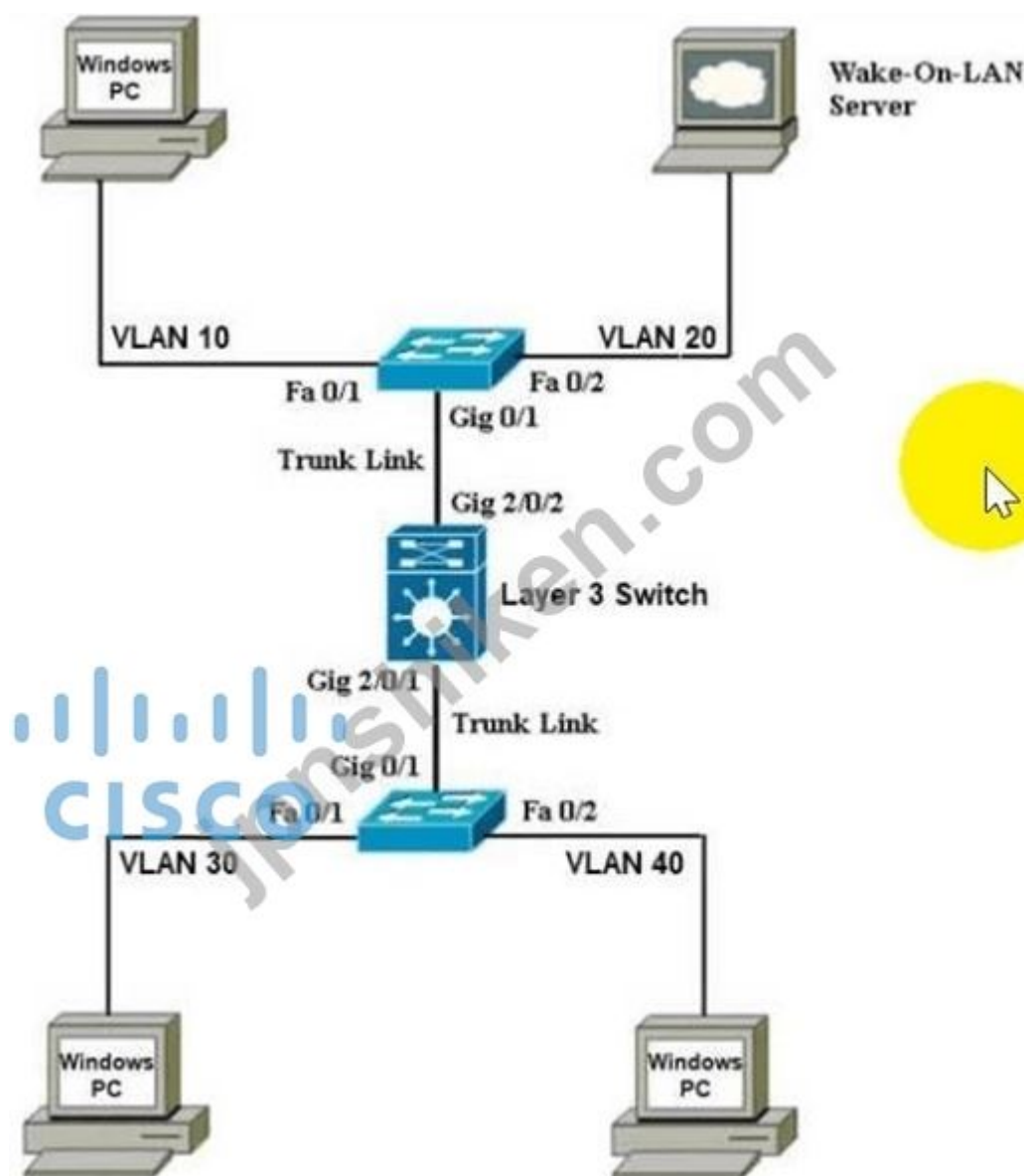
正解: ([正解を表示します](#))

StackWise is the correct solution because the company needs more access ports while reusing the existing access-switch uplink capacity. Cisco StackWise allows multiple compatible Catalyst switches to operate as a single logical switching system, increasing access-port density without requiring each added switch to consume new uplink ports on the distribution layer. This fits the design constraint that the distribution switches have no available ports for additional uplinks, while current uplink bandwidth is not a bottleneck.

VSS is normally used to combine two chassis-class switches into one logical system at the distribution or core layer; it is not the practical answer for expanding one floor's access-port count from an existing access switch.

EtherChannel aggregates physical links for bandwidth and redundancy, but it does not create additional access ports when no distribution ports are available. VDC is a Nexus virtualization construct and does not solve Catalyst access-port expansion. The correct operational design is to add a stack member and use the existing uplink design. Reference topics: Cisco StackWise, access-layer scalability, uplink preservation, Catalyst campus switching.

質問: 162



展示を参照してください。従業員 ID: 1234 56:789 のエンジニアは、クライアント用の WoL 展開を設計する必要があり、その設計では、サーバー側が命令を開始したときに Windows PC が WoL マジック パケットに遅延なく応答することを保証する必要があります。エンジニアはどのアクションを選択する必要がありますか？

- A. クライアントが存在するすべてのインターフェイスでスパンニングツリー PortFast を有効にする必要があります。
- B. WoL はネットワーク カードで有効にし、Windows PC の BIOS で無効にする必要があります。
- C. クライアントが存在するすべてのインターフェイスで IP ダイレクト ブroadcast を無効にする必要があります。
- D. クライアントが存在するすべてのインターフェイスで IP 転送プロトコルを無効にする必要があります

正解: (正解を表示します)

Spanning-tree PortFast must be enabled on interfaces where Wake-on-LAN clients reside so the PCs become reachable immediately when they power on. WoL depends on a sleeping endpoint receiving a magic packet, and once the device wakes, it must be able to send and receive traffic without waiting through traditional STP listening and learning states. Cisco campus designs commonly enable PortFast on access ports connected to end hosts so those ports transition directly to forwarding. That reduces login, DHCP, and application delays after a host becomes active. The NIC and system firmware must support WoL, but the option stating that WoL should be disabled in BIOS is technically wrong. IP directed broadcast and UDP forwarding are relevant when WoL packets must cross routed boundaries, but the answer options here focus on a single action to prevent client-side delay. Disabling directed broadcast or forward protocol would work against remote WoL operation. Reference topics: Wake-on-LAN, PortFast, access-layer design, STP edge ports, client startup convergence.

質問: 163

ファブリック データ プレーンは SD-Access アーキテクチャで何を利用しますか？

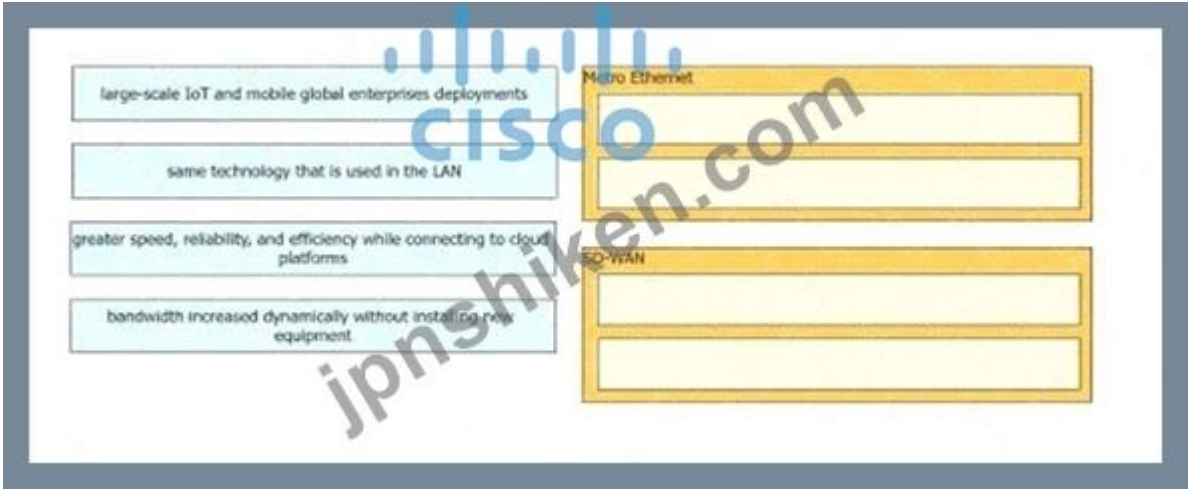
- A. エンドポイントとロケーションのマッピングを解決するための LISP プロトコル
- B. リンクステート ルーティング情報を交換する IS-IS プロトコル
- C. レイヤ 2 フレームを転送するための MAC-in-IP カプセル化方式
- D. ファブリックの外部にエンドポイント プレフィックスをアドバタイズするための BGP プロトコル

正解: [\(正解を表示します\)](#)

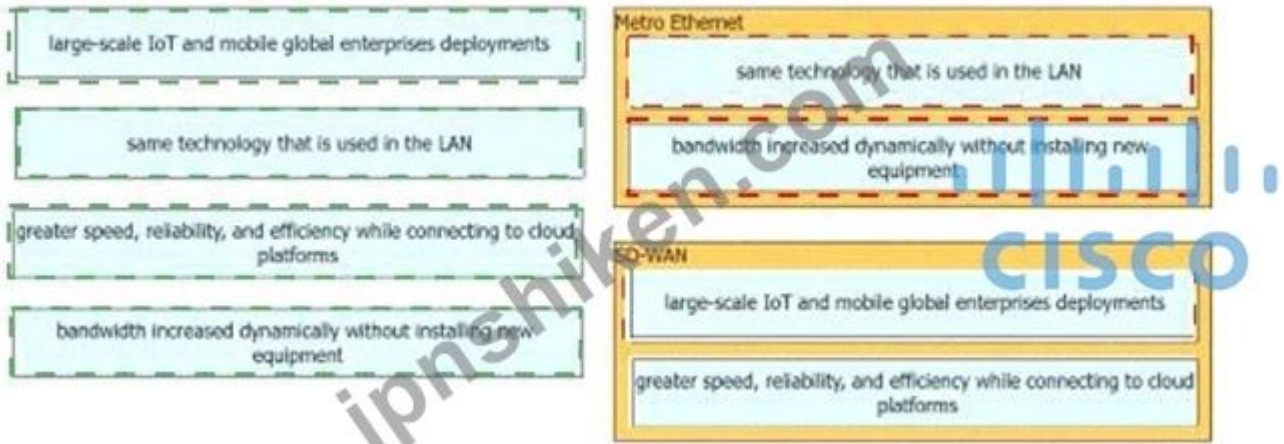
The Cisco SD-Access fabric data plane leverages VXLAN-style MAC-in-IP encapsulation to transport endpoint traffic across the Layer 3 fabric. In SD-Access, the underlay provides IP reachability between fabric nodes, while the overlay uses encapsulation to carry user traffic and policy context across that underlay. Cisco fabric edge nodes encapsulate traffic before forwarding it through intermediate nodes and decapsulate it when it reaches the destination edge or border. LISP is used by the control plane to resolve endpoint-to-location mappings, so option A describes the mapping function rather than the data plane. IS-IS provides underlay routing reachability and is not the overlay data-plane mechanism. BGP can be used outside the fabric or in other overlay designs, but it is not the SD-Access fabric data plane. The exam option describes the essential encapsulation function, even though VXLAN is technically MAC-in-UDP over IP. Reference topics: SD- Access data plane, VXLAN encapsulation, fabric edge forwarding, underlay and overlay separation, policy transport.

質問: 164

左側の説明を右側の該当するカテゴリにドラッグ アンド ドロップします。



正解:



Explanation:

The drag-and-drop mapping should be preserved as the validated answer because the item depends on the exhibit categories and the original coordinate map identifies which descriptions belong to each category. In Cisco design questions of this form, the important skill is distinguishing the functional category from the implementation detail. A professional design answer must map each description to the correct operational class instead of treating all items as interchangeable features. This is common in ENSLD topics such as WAN connectivity, SD-Access roles, telemetry modes, and QoS service models,

where a description may sound similar but belongs to a different architectural layer. The coordinate answer keeps the specific pairings intact while the explanation documents the design logic: categorize by function, control-plane responsibility, data- plane behavior, and operational use case. If the exhibit contains technology categories, the selected mapping should be read as the official answer for those categories rather than converted into free text. Reference topics: Cisco design categorization, architectural roles, technology-to-function mapping, ENSLD drag-and- drop interpretation.

質問: 165

左側の特性を、右側の適用先のテレメトリ モードにドラッグ アンド ドロップします。

The collector initiates a session to the device.

supports TCP, UDP, and gRPC.

The device initiates a session to the collector.

supports gRPC only.

Dial-In

Dial-Out

正解:

The collector initiates a session to the device.

supports TCP, UDP, and gRPC.

The device initiates a session to the collector.

supports gRPC only.

Dial-In

The collector initiates a session to the device.

supports gRPC only.

Dial-Out

The device initiates a session to the collector.

supports TCP, UDP, and gRPC.

Explanation:

Dial-In: collector initiates a session to the device; supports gRPC only. Dial-Out: device initiates a session to the collector; supports TCP, UDP, and gRPC.

Model-driven telemetry can be deployed in dial-in or dial-out mode, and the difference is which side initiates the session and where the subscription is controlled. In dial-in mode, the collector initiates the session to the network device and requests the telemetry subscription. This makes the collector responsible for establishing the session and asking for the data stream. In the context of the provided options, dial-in is paired with support for gRPC only. In dial-out mode, the network device initiates the session toward the collector according to a configured subscription. This is useful when devices sit behind NAT or when the operational model requires devices to push telemetry to a known collector. In the provided choices, dial-out is paired with support for TCP, UDP, and gRPC. The mapping is not arbitrary; it follows Cisco telemetry terminology, where dial-in is collector-initiated and dial-out is device-initiated. This distinction is important in network automation designs because firewall policy, collector scaling, and subscription lifecycle management differ between the two modes.

質問: 166

建築家は、会社のメインサイトをいくつかの中小規模のリモートブランチに接続するための設計に取り組んでいます。ソリューションには冗長WANリンクを含める必要がありますが、お客様の予算は限られており、将来的にリンク速度を簡単に上げることができることを望んでいます。QoSはブランチルーターにはないため、一貫したエンドツーエンドのQoSは必要ありません。アーキテクトはどのソリューションを提案しますか？

- A. シングルエッジルーターを備えたデュアルホームWAN MPLS
- B. サイト間VPNトポロジを実行するシングルエッジルーターを備えたデュアルホームインターネット

C. デュアルエッジルーターを介したデュアルホームWANMPLSおよびインターネットリンク

D. ハブアンドスポークVPNトポロジを実行するデュアルエッジルーターを備えたデュアルホームインターネット

正解: ([正解を表示します](#))

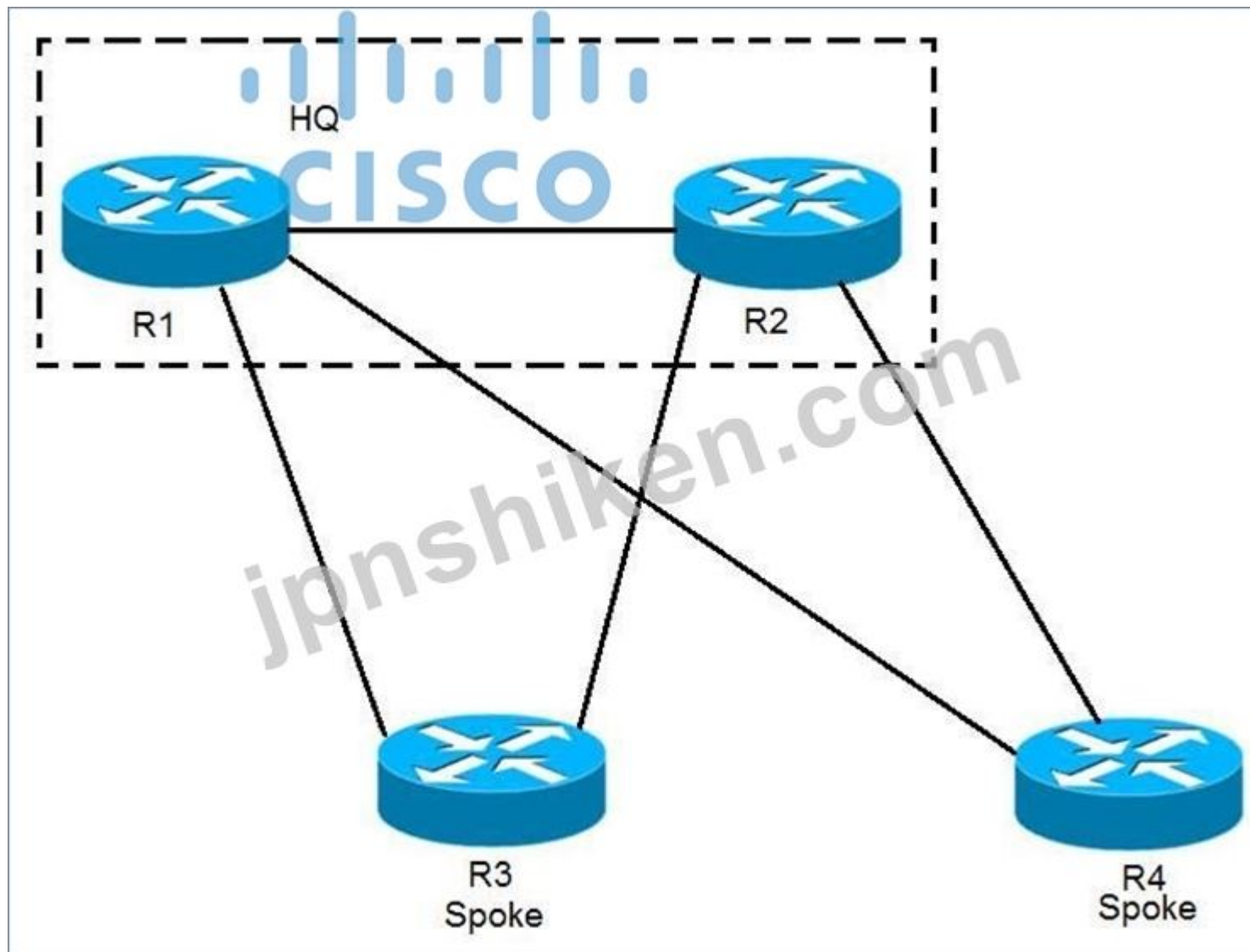
A dual-homed Internet design with a single edge router and site-to-site VPN topology best matches the stated cost and flexibility requirements. The customer needs redundant WAN links, has a limited budget, wants easy future bandwidth increases, and does not require consistent end-to-end QoS on the branch routers. Internet circuits are typically less expensive and easier to upgrade than MPLS access circuits, making them appropriate for small to medium-sized remote branches when strong QoS guarantees are not required. A single edge router also reduces hardware cost and operational complexity compared with dual edge routers.

WAN MPLS designs can provide stronger service-level guarantees and QoS treatment, but they are more expensive and less flexible for rapid bandwidth upgrades. A design with dual edge routers improves device redundancy but conflicts with the limited-budget condition. Therefore, a dual-homed Internet solution using a site-to-site VPN overlay is the most appropriate option. The architect should still validate encryption requirements, failover behavior, routing preference, and whether the Internet underlay can meet application performance expectations. Reference topics: branch WAN design, dual Internet connectivity, site-to-site VPN, cost-driven WAN selection.

有効的な**300-420J**問題集はJPNTTest.com提供され、**300-420J**試験に合格することに役に立ちます！JPNTTest.comは今最新**300-420J**試験問題集を提供します。JPNTTest.com 300-420J試験問題集はもう更新されました。ここで**300-420J**問題集のテストエンジンを手に入れます。最新版のアクセス、<https://www.jpntest.com/shiken/300-420J-mondaishu> **457**問、**30%ディスカウント**、特別な割引コード：**JPNshiken**」

質問: **167**

展示品を参照してください。



すべてのリンクに EIGRP が設定されています。スポーク ノードは EIGRP スタブとして設定されており、R3 への WAN リンクは R4 へのリンクよりも帯域幅が高く、遅延が低くなっています。R1-R2 リンクでリンク障害が発生した場合、R2 に接続されたサブネット宛ての R1 上のトラフィックはどうなりますか？

- A. R1はR2へのルートを持たず、トラフィックをドロップします。
- B. R1はR3とR4を経由してR2に到達するパスに負荷分散します。
- C. R1はトラフィックをR3に転送しますが、R3はトラフィックをドロップします。
- D. R1はR2に到達するためにトラフィックをR3に転送します。

正解: ([正解を表示します](#))

EIGRP stub routing is designed to prevent remote or spoke routers from being used as transit routers. In a hub- and-spoke or dual-homed remote design, a stub router advertises only the route types permitted by its stub configuration, such as connected or summary routes, and it is not queried for routes that it cannot reasonably provide. This reduces EIGRP query scope and helps avoid stuck-in-active conditions. The tradeoff is that the network must not depend on a stub site to carry transit traffic between other parts of the topology. In this scenario, the spoke nodes are configured as EIGRP stubs. When the direct R1-to-R2 link fails, R1 cannot use the alternate path through the other spoke-side topology as a transit path to reach a subnet attached to R2. The EIGRP stub design prevents the route from being advertised in a way that would make the spoke act as a transit router. As a result, R1 has no valid route to R2 's attached subnet and drops the traffic. This behavior is consistent with the purpose of EIGRP stub routing.

質問: 168

マルチキャスト ネットワークでは、RP アドレスに対して RPF チェックを実行するためにどの条件を満たす必要がありますか？

- A. PIM DMデバイスはマルチキャストパケットを受信し、直接接続されたメンバーが存在しない

B. PIMルータまたはマルチレイヤスイッチは共有ツリー状態です

C. PIMルータまたはマルチレイヤスイッチはソースツリー状態です

D. PIM DMデバイスはマルチキャストパケットを受信し、直接接続されたPIMネイバーがありません。

正解: ([正解を表示します](#))

The RPF check is performed against the RP address when the PIM router has shared-tree state. In PIM sparse mode, multicast traffic initially flows through the rendezvous point using the shared tree, represented as (*,G) state. For shared-tree traffic, the router verifies that packets arrive from the direction of the RP, because the RP is the root of the shared distribution tree. When the router has source-tree state, represented as (S,G), the RPF check is performed against the multicast source address instead. Dense mode behavior and directly connected membership are not the determining conditions in this question. The key distinction is whether forwarding state is shared-tree or source-tree. Cisco multicast design relies on this behavior to keep PIM forwarding loop free as traffic transitions from shared-tree delivery to shortest-path tree delivery where applicable. Selecting the shared-tree state option correctly identifies when the RP address is used for the RPF calculation. Reference topics: PIM sparse mode, shared tree, source tree, RPF toward RP, (*,G) and (S,G) state.

有効的な**300-420J**問題集はJPNTTest.com提供され、**300-420J**試験に合格することに役に立ちます！JPNTTest.comは今最新**300-420J**試験問題集を提供します。JPNTTest.com 300-420J試験問題集はもう更新されました。ここで**300-420J**問題集のテストエンジンを手に入れます。最新版のアクセス、<https://www.jpntest.com/shiken/300-420J-mondaishu> **457**問、**30%ディスカウント**、特別な割引コード：**JPNshiken**」