

HCIE-Datacom V1.0

Huawei H12-891 V1.0

Version Demo

Total Demo Questions: 63

Total Premium Questions: 633

Buy Premium PDF

<https://passqueen.com>

support@passqueen.com

Topic Break Down

Topic	No. of Questions
Topic 1, Datacom Network Architecture	382
Topic 2, Datacom Network Planning and Design	21
Topic 3, Datacom Network Implementation and O&M	145
Topic 4, Datacom Network Optimization	54
Topic 5, Datacom Network Troubleshooting	31
Total	633

QUESTION NO: 1

The following description of the RD attribute filter for BGP is correct?

- A. If RD-filter is not configured, but the RD-filter is referenced for filtering, the matching result is deny.
- B. If RD-filter is configured, but the routed RD is not matched with any of the RDs defined in the rule, the default match result is permit.
- C. Multiple rules are matched in the order in which they are configured.
- D. There is always an "or" relationship between the rules configured by RD-filter.

ANSWER: A C D

Explanation:

An RD filter is evaluated as an ordered set of rules when it is referenced by BGP VPN route filtering. Therefore, "Multiple rules are matched in the order in which they are configured." is correct: the device evaluates entries sequentially and uses the applicable result according to that configured sequence. "There is always an "or" relationship between the rules configured by RD-filter." is also correct because an RD only needs to match one applicable filter rule; the rules are not combined as cumulative AND conditions.

"If RD-filter is not configured, but the RD-filter is referenced for filtering, the matching result is deny." is correct as well. A referenced but nonexistent RD filter cannot provide a permitted match, so the route fails the match and is denied. This behavior prevents unintentional acceptance of VPN routes when a filtering object has not been created. Route distinguishers are the values used to make otherwise overlapping VPN-IPv4 or VPN-IPv6 prefixes unique, as defined by [RFC 4364](#). Huawei's enterprise documentation portal provides the VRP BGP VPN and routing-policy configuration references for these filter-application semantics: [Huawei Enterprise Support](#).

QUESTION NO: 2

There are many ways to establish a cluster on a box switch, what methods does Huawei frame switch support to establish a cluster?

- A. Set up via a cluster card on the motherboard
- B. Establish
- C. through the cluster daughter card on the interface board. Establish
- D. through the service interface on the interface board Established by switching the cluster card on the network board

ANSWER: A C D

Explanation:

Huawei chassis switches support Cluster Switching System connections through dedicated cluster hardware and, on supported models, through service interfaces. A cluster card installed on the main control board provides dedicated physical resources for inter-chassis cluster links. This approach is intended for chassis platforms that provide cluster-card slots on their main control boards. A cluster daughter card installed on an interface board is another supported hardware-based implementation, allowing the interface-board architecture to provide the cluster connection. In addition, supported chassis can establish cluster links by using service interfaces on interface boards or dedicated cluster cards on switching-network boards. These alternatives let the platform use the card and port resources available in its particular hardware design while forming the logical cluster and forwarding traffic as one system. The exact available connection method depends on the chassis model, installed boards, and software version; the hardware installation guide and CSS configuration documentation for the specific device should therefore be checked before deployment. Huawei's enterprise support portal provides the applicable product documentation at [Huawei Enterprise Support](#), and Huawei's switch product documentation is available from [Huawei Switches Documentation](#).

QUESTION NO: 3

Link state information is isolated by default between different IS-IS processes on the same router.

- A. TRUE
- B. FALSE

ANSWER: A

Explanation:

TRUE is correct. An IS-IS process is a separate routing-protocol instance with its own protocol configuration, neighbor adjacencies, link-state database, shortest-path-first calculations, and routing information generated from that instance. When multiple IS-IS processes are configured on one Huawei router, each process is identified separately and operates independently. Consequently, Link State PDUs received and processed in one process populate only that process's link-state database; they are not automatically shared with another local IS-IS process.

This isolation is fundamental to running separate IS-IS routing domains, such as distinct service, tenant, or migration networks, on the same device. Any intended exchange of reachability information between independently configured routing processes must be configured explicitly through route import or redistribution policy. It is not an inherent sharing of IS-IS link-state information. Huawei's IS-IS configuration documentation describes creating IS-IS processes as separate instances and configuring interfaces to participate in the required process. The protocol behavior and link-state database model are also defined by the IS-IS standard, in which a router maintains the database associated with each routing-domain instance. See Huawei's [IS-IS configuration documentation](#) and [ISO/IEC 10589 IS-IS specification](#).

QUESTION NO: 4

::1/128 is the IPV6 loopback address

- A. TRUE
- B. False

ANSWER: A

Explanation:

TRUE is correct because ::1/128 is the IPv6 loopback address. The address is written in full as 0:0:0:0:0:0:0:1, with IPv6 zero-compression reducing the consecutive all-zero hextets to ::. The /128 prefix length identifies one individual IPv6 address, rather than a subnet or address range. A host uses the loopback address to send packets to its own IPv6 protocol stack. This is useful for local testing and for verifying that IPv6 is enabled and operating on the device, independently of any physical interface or external network path. For example, a ping to ::1 should be processed locally and must not be forwarded onto a network link. The IPv6 addressing architecture reserves this address specifically for loopback operation and states that a packet with a loopback source or destination address must not be sent beyond the originating node. Huawei IPv6 documentation likewise uses ::1 as the standard IPv6 loopback address. See [RFC 4291, Section 2.5.3](#) and [Huawei IPv6 Addresses documentation](#).

QUESTION NO: 5

Regarding the difference between NSR and NSF, the correct one is?

- A. The NSR must rely on a neighbor router to complete
- B. Both NSR and NSF require a neighbor router to complete
- C. NSF must rely on neighbor routers to complete
- D. NSF does not require a neighbor router to complete

ANSWER: C

Explanation:

NSF must rely on neighbor routers to complete. Nonstop Forwarding is a graceful-restart mechanism intended to keep the data plane forwarding packets while a routing protocol process or control plane undergoes a restart. To avoid an immediate loss of routes during that interval, the restarting router signals its neighbors, and capable neighbors act as graceful-restart helpers. They retain the restarting router's routing information for a negotiated restart period and continue forwarding traffic based on the existing topology until the router recovers and re-establishes protocol state. Therefore, NSF is not achieved by the restarting device in isolation: neighboring routers must support and participate in the relevant protocol's graceful-restart/helper behavior. The exact signaling and retained-state behavior are protocol-specific. For example, OSPF graceful restart defines helper operation and the conditions under which neighbors maintain adjacencies during a restart, while BGP graceful restart defines the exchange of restart capability and the retention of forwarding state. These standards describe the inter-router cooperation that makes NSF possible. See the IETF's [OSPF Graceful Restart specification](#) and [BGP Graceful Restart specification](#).

QUESTION NO: 6

Regarding the DSCP field, the following statement is incorrect: single selection

- A. The DSCP value can be 0, and packets use the default forwarding mechanism when 0
- B. When the value of DSCP is EF, it indicates that the data belongs to the accelerated forwarding class.
- C. In addition to indicating the priority of the data, DSCP can also indicate the probability of data being discarded. The CS class in D. DSCP divides data into 8 priorities.

ANSWER: B

Explanation:

"When the value of DSCP is EF, it indicates that the data belongs to the accelerated forwarding class." is the incorrect statement because EF is the standardized DiffServ Per-Hop Behavior name for *Expedited Forwarding*, not "accelerated forwarding." EF is intended for traffic requiring very low loss, low latency, low jitter, and assured bandwidth, so that it receives a forwarding treatment comparable to a virtual leased line. The EF DSCP codepoint is 101110 (decimal 46), and network devices can use this marking to classify traffic into an appropriate high-priority QoS behavior. Correct terminology matters in Huawei QoS configuration and troubleshooting because classifiers, traffic behaviors, queues, and policies are based on defined DiffServ service classes and PHB semantics. RFC 3246 formally defines the Expedited Forwarding PHB and its operational intent, while RFC 2474 defines the DS field and the DSCP structure used for differentiated services. See [RFC 3246](#) and [RFC 2474](#).

QUESTION NO: 7

Mainstream Layer 2 technologies such as VXLAN, TRILL, NVGRE and MPLS

- A. True
- B. False

ANSWER: A

Explanation:

True is correct. VXLAN, TRILL, NVGRE, and MPLS are all technologies commonly encountered when designing or providing Layer 2 connectivity, particularly in data-center and service-provider environments. VXLAN encapsulates Ethernet frames over an IP underlay and uses a 24-bit VXLAN Network Identifier, allowing Layer 2 segments to be extended across a routed Layer 3 network at large scale. The VXLAN encapsulation method is standardized in [RFC 7348](#). TRILL was designed to improve Layer 2 forwarding in Ethernet networks by using routing principles and encapsulating frames for forwarding across a TRILL campus; its base protocol is specified in [RFC 6325](#). NVGRE similarly transports tenant Ethernet traffic over an IP network using GRE encapsulation. MPLS is often described as a Layer 2.5 forwarding technology because it forwards

labeled packets between Layer 2 and Layer 3 processing, but it is extensively used to deliver Layer 2 VPN services such as VPWS and VPLS. Therefore, in the practical context of technologies used to implement scalable Layer 2 networks and Layer 2 services, the listed technologies are appropriately grouped together.

QUESTION NO: 8

If the Option field value in the OSPFv3 Hello message issued by the router's Gigabit Ethernet0/0 interface is 0x000013, the following description is correct

- A. The Gigabit Ethernet0/0/0 interface of Router A adds IPv6 routing to compute
- B. Router A's Gigabit Ethernet0/0/0 interface belongs to the NSSA zone
- C. Router A supports AS-External-LSA flood
- D. Router A is an OSPFv3 device with forwarding capability

ANSWER: A C D

Explanation:

The hexadecimal Options value 0x000013 has the V6, E, and R capability bits set: $0x01 + 0x02 + 0x10 = 0x13$. The V6 bit indicates that the router supports IPv6 forwarding and participates in IPv6 routing functionality, which supports the statement that the interface adds IPv6 routes for computation. The E bit advertises external-routing capability. Consequently, the router supports flooding AS-External-LSAs in the relevant OSPFv3 area scope. The R bit means that the router is active for routing and is capable of forwarding traffic; therefore, it is an OSPFv3 router with forwarding capability. These Options bits are carried in OSPFv3 Hello packets so neighboring routers can understand the capabilities applicable to the adjacency and area. The Network bit, which denotes NSSA capability, is not present in 0x13; the set bits are specifically V6, E, and R. The definitions and bit assignments are specified in the OSPFv3 protocol specification and remain the basis for interpreting Hello-packet Options fields on Huawei routers. See [RFC 5340, OSPF for IPv6](#) and Huawei's [OSPF configuration documentation](#).

QUESTION NO: 9

What are the types of link types included in OSPF's route LSA?

- A. Stubnet
- B. P-2-P
- C. VLINK
- D. Transit network

ANSWER: A B C D

Explanation:

An OSPF Router-LSA (Type 1 LSA), sometimes called a route LSA in training material, describes the router's links within its area. The Router-LSA link-description format defines exactly four link types: a stub network, a point-to-point connection to another router, a connection to a transit network, and a virtual link. Therefore, *Stubnet*, *P-2-P*, *VLINK*, and *Transit network* are all valid link types represented in a Router-LSA. Each link entry includes the link ID, link data, link type, metric, and optional type-of-service information, allowing SPF calculation to construct the area topology accurately. Point-to-point and virtual links describe router adjacencies, while stub and transit-network links represent reachable network segments advertised by the router. RFC 2328 specifies these Router-LSA link types and their numeric encodings: point-to-point is 1, transit network is 2, stub network is 3, and virtual link is 4. See the Router-LSA definition in [RFC 2328, section 12.4.1](#) and the link-type field details in [RFC 2328, Appendix A.4.2](#).

QUESTION NO: 10

Regarding the MPLS handling mode of TTL, the following description is correct?

- A. Pipe mode. When an IP message passes through an MPLS network, the IPTTL minus 1 at the entry point is mapped to the MPLS TTL field
- B. The TTL in an MPLS label has the same purpose as the TTL field in an IP header: it prevents MPLS routing loops.
- C. In Uniform mode, when an IP message passes through an MPLS network, the inbound IPTTL is minus 1. The MPLSTTL field is a fixed value
- D. In MMLSVP, if you need to hide the structure of the MPLS backbone network, you can use the Uniform mode on Ingress for private network messages

ANSWER: B

Explanation:

The correct statement is: "The TTL in an MPLS label has the same purpose as the TTL field in an IP header: it prevents MPLS routing loops." MPLS carries an 8-bit TTL field in every label-stack entry. Its fundamental function is equivalent to the IP header TTL (or IPv6 Hop Limit): it limits the number of forwarding hops a packet can take. At each MPLS-capable forwarding hop, the label TTL is decremented. If it reaches zero, the forwarding device discards the packet and normally generates an ICMP Time Exceeded message where applicable. This bounded lifetime prevents a packet from circulating indefinitely when a forwarding or label-switched-path loop exists, protecting network resources and making loop faults observable through traceroute-style diagnostics. The MPLS architecture defines TTL processing as part of the label-stack encoding, and MPLS TTL-processing procedures preserve this loop-prevention role regardless of whether the network uses Uniform, Pipe, or Short Pipe TTL behavior. Those modes chiefly control how IP-layer and MPLS-layer TTL values are propagated at LSP boundaries; they do not change the essential anti-loop purpose of the MPLS label TTL. See [RFC 3032, MPLS Label Stack Encoding](#) and [RFC 3443, TTL Processing in MPLS Networks](#).

QUESTION NO: 11

For Layer 2 VPN technology, the following statement is correct? Multi-select

- A. VPLS is a widely used technology in the live network, which can provide transparent transmission of three-layer messages and achieve multi-point access.
- B. The VPLS configuration is complex and the Layer 2 network transmits BUM packets.
- C. BGPEVPN supports features such as tenant isolation, Multi-homing, and broadcast suppression. D. BGPEVPN addresses mac address drift and multi-tenancy that VPLS cannot support.

ANSWER: B C

Explanation:

"The VPLS configuration is complex and the Layer 2 network transmits BUM packets." is correct because VPLS creates a multipoint Ethernet service across an MPLS/IP provider network. To emulate a shared LAN, a VPLS PE must learn customer MAC addresses and distribute broadcast, unknown-unicast, and multicast (BUM) traffic to the relevant remote PEs. This flooding behavior, together with the required signaling, pseudowire, and split-horizon considerations, can make VPLS operationally complex as the service scales. The VPLS architecture and its multipoint forwarding behavior are described in [RFC 4761](#).

"BGPEVPN supports features such as tenant isolation, Multi-homing, and broadcast suppression. D. BGPEVPN addresses mac address drift and multi-tenancy that VPLS cannot support." is also correct in substance. BGP EVPN uses BGP control-plane routes to advertise MAC/IP reachability, enabling more controlled forwarding than data-plane MAC learning and flooding alone. EVPN supports Ethernet segment multihoming, tenant separation through VPN instances and route targets, MAC mobility handling when endpoints move, and mechanisms that reduce unnecessary BUM replication. These capabilities make EVPN especially suitable for scalable multi-tenant data-center and enterprise VPN deployments. The EVPN control-plane model and MAC mobility procedures are specified in [RFC 7432](#).

QUESTION NO: 12

In an MPLS IBGP VPN environment, if tags are distributed only through BGP, LDP, the MPLS tags for a message can be up to two layers of tags.

- A. True
- B. False

ANSWER: A

Explanation:

True is correct. In a conventional MPLS Layer 3 VPN deployment using internal BGP for VPN route advertisement, a customer packet normally carries a two-label MPLS stack while it traverses the provider backbone. The inner VPN label is advertised with the VPN route by Multiprotocol BGP and tells the egress provider edge device which VPN forwarding context and outbound route to use. The outer transport label is learned through LDP and carries the packet across the provider MPLS core toward the egress provider edge device. Each transit provider router processes only the outer transport label, swapping it as needed; the VPN label remains until the packet reaches the egress edge. The outer label is then removed, allowing the egress device to use the remaining VPN label for customer-VPN forwarding. Therefore, when the stated environment relies on BGP for the VPN label and LDP for the backbone transport label, the packet has no more than these two MPLS label layers. This behavior follows the BGP/MPLS IP VPN forwarding model specified in [RFC 4364](#) and the MPLS label-stack architecture described in [RFC 3031](#).

QUESTION NO: 13

Which of the following scenarios send free ARP messages?(Multiple choice questions).

- A. After the VRRP protocol master and standby role election is completed
- B. The DHCP client receives an acknowledgment message from the server after it sends it
- C. When a new host configured with an IP address is connected to the network
- D. When a host connected to the switch issues a Ping message.

ANSWER: A B C

Explanation:

Gratuitous ARP (also called free ARP) is an unsolicited ARP announcement sent using the sender's own IPv4 address as the target protocol address. Its purpose is to inform other devices of an IP-to-MAC mapping change, refresh ARP caches, and help detect duplicate IPv4 address use. After VRRP master and standby role election is completed, the device that becomes Master sends gratuitous ARP packets for the virtual IPv4 address. This promptly updates connected hosts and switches so traffic for the virtual gateway is forwarded to the current Master. The VRRP specification describes the Master's responsibility to send gratuitous ARP when assuming the Master role: [RFC 3768](#).

The DHCP client receives an acknowledgment message from the server after it sends it is also a valid trigger in common DHCP operation: after receiving a DHCPACK and before using the assigned address, a client can use ARP to check whether the address is already in use and announce its ownership when configuring the address. DHCP address-conflict checking through ARP is specified in [RFC 2131](#). Similarly, when a new host configured with an IP address is connected to the network, it commonly sends gratuitous ARP to advertise its newly active IP-to-MAC binding and detect any duplicate address condition. These behaviors enable neighboring devices to learn the current mapping without waiting for a normal ARP resolution exchange.

QUESTION NO: 14

Two routers connected through a firewall cannot establish a BGP neighbor relationship. What causes it?

- A. EBGp Multi-hop is not configured

- B. The firewall did not release IP protocol number 89
- C. Firewalls that do not have an open TCP port number of 179
- D. BGP neighbors must form a Layer 2 adjacency

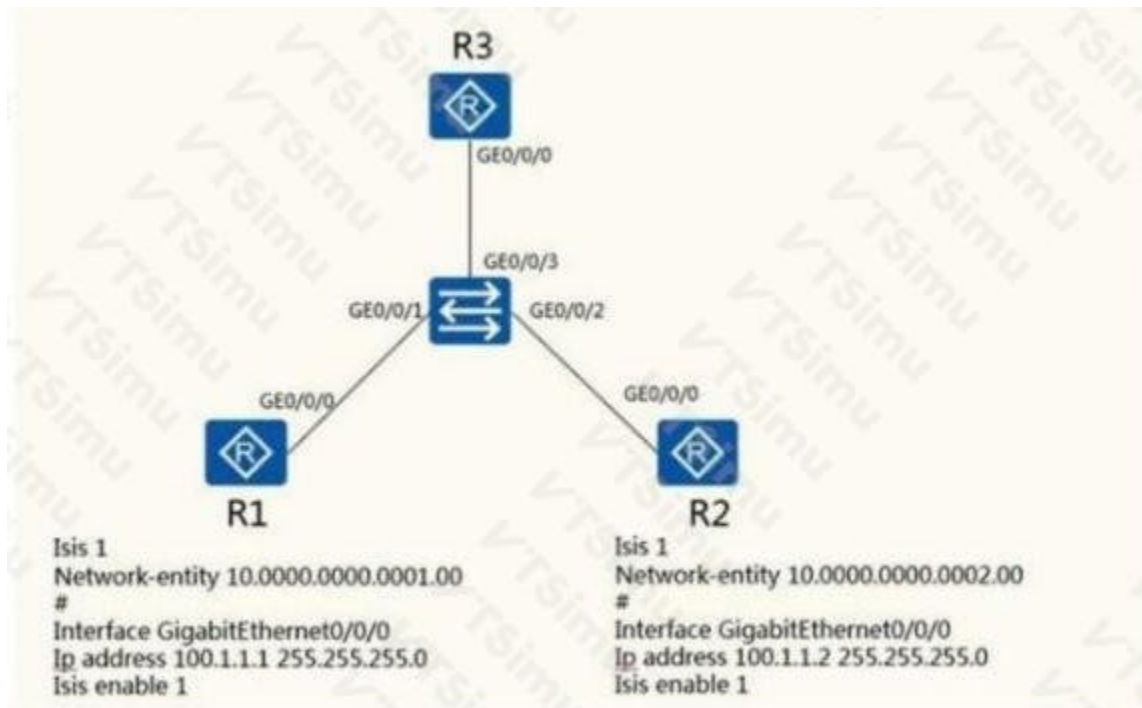
ANSWER: C

Explanation:

Firewalls that do not have an open TCP port number of 179 is correct because BGP establishes its neighbor session using TCP, with TCP destination port 179. Before BGP can exchange OPEN messages, keepalives, updates, or notifications, one peer must successfully create the underlying TCP connection to the other peer. A firewall positioned between the routers must therefore permit TCP/179 in the required direction and allow the corresponding return traffic. For a stateful firewall, this normally means an interzone security policy permitting a new TCP session to destination port 179; return packets are then permitted as part of the established session. If TCP/179 is filtered, the TCP three-way handshake cannot complete, so the BGP finite-state machine cannot advance to an established neighbor relationship. This requirement applies whether the peers are directly connected or are reachable across an intervening routed network. The BGP specification explicitly defines TCP port 179 as the well-known port used by BGP speakers, and Huawei's [BGP configuration documentation](#). See [RFC 4271, section 8.1](#) and Huawei's [BGP configuration documentation](#).

QUESTION NO: 15

As shown in the figure, the overload bit of ISIS is set in the configuration of router R3, and the following statement is correct



- A. R3's direct route will be ignored by R1 and R2 because of the setting of the overload bit .
- B. When the router runs out of memory, the system automatically sets the overload bit in the sent LSP packets, regardless of whether the user has configured the set-overload command.
- C. Routes introduced by R3 through import-route will not be published in R3's LSPs.
- D. The LSP of R3 will not be able to spread to R1 and R2.
- E. R3 initiates the overload state, and R1 and R2 will not calculate R3's LSP in the SPF tree.

ANSWER: B C

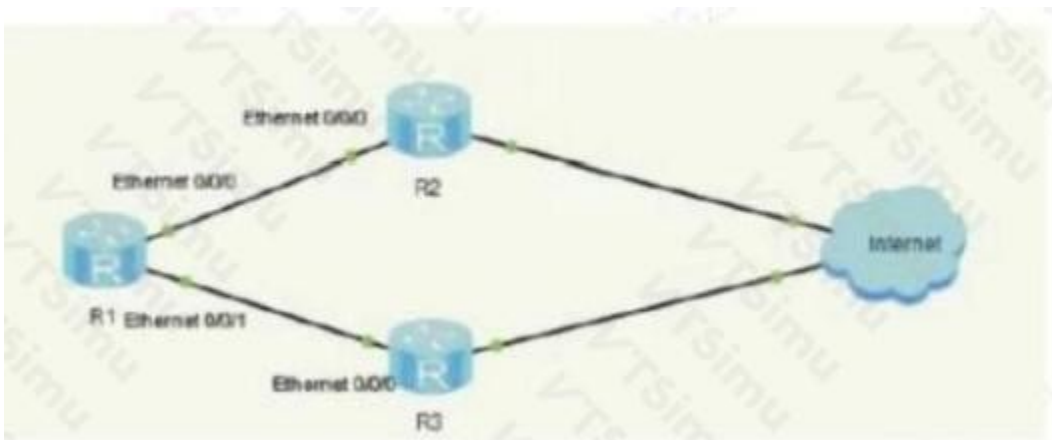
Explanation:

When an IS-IS router has insufficient memory to maintain the required link-state information, it can automatically set the overload bit in its generated LSPs. This protection is independent of a manually configured `set-overload` command: manual configuration is an administrative way to place the router in overload state, whereas a resource shortage can cause the protocol to assert the bit automatically. The overload indication enables surrounding routers to avoid relying on an impaired router for normal transit forwarding while the router recovers or is being maintained.

With R3 in overload state, routes that R3 imports from another routing protocol by using `import-route` are not advertised in R3's IS-IS LSPs. This behavior prevents redistributed reachability from being propagated through a router that has declared itself overloaded. It is especially useful when controlling traffic during startup, maintenance, or resource exhaustion, because it limits the risk that externally learned destinations are advertised through an unstable transit point. The IS-IS overload-bit mechanism and its use in LSP processing are specified in [RFC 1195](#); Huawei configuration documentation and command references are available through the [Huawei Enterprise Support portal](#).

QUESTION NO: 16

In the following topology diagram,



R1 can be established between R2 and R3 fetches the Internet, R1 and R2, R1 and R3 EGP connections, R2, R3 belong to the same AS. The default routes are issued through BGP separately, requiring R2 to be preferred for traffic accessed by R1 to the Internet, and the following practices can be achieved

- A. R1 specifies that peer R3 sets the MED property to 150 in the import direction
- B. R2 specifies that peer R1 is in export. The direction setting local priority is 150
- C. R2 specifies peer R1 in the export direction setting MED property to 150
- D. R1 specifies the peer R2 sets the local priority in the import direction to 150

ANSWER: A D

Explanation:

R1 can prefer the Internet default route through R2 by assigning a higher local preference to the default route learned from R2 in the inbound direction. Local preference is a BGP attribute used inside an AS, and BGP selects the route with the higher local-preference value. Therefore, assigning a value of 150 to the routes received from R2 makes that path preferred over the corresponding route received from R3, assuming the latter retains the normal lower value. This is the most direct policy method because the decision is made locally on R1 before the route is installed and used for forwarding.

R1 can also set MED to 150 on the route received from R3. R2 and R3 are in the same neighboring AS, so their routes are eligible for normal MED comparison. BGP prefers the path with the lower MED value; consequently, increasing the MED of the default route learned from R3 causes the default route learned through R2 to be selected, provided its MED remains

lower. Huawei routing policies can modify BGP attributes in the inbound direction, allowing either approach to influence R1's best-path calculation. See the BGP attribute definitions in [RFC 4271](#) and Huawei's [BGP configuration documentation](#).

QUESTION NO: 17

The network administrator is ready to use the 6to4 automatic tunnel to carry IPv6 data on the IPv4 network, and the IPv4 address of the router interface is 1381485210, so the corresponding tunnel address is 2001:8a0e:55d2:1:230:65ff:fe2:9a6?

- A. True
- B. False

ANSWER: B

Explanation:

False is correct because a 6to4 IPv6 address is formed from the dedicated 6to4 prefix 2002::/16, followed by the 32-bit IPv4 address represented in hexadecimal. The decimal IPv4 value 1381485210, interpreted as an unsigned 32-bit IPv4 address, converts to hexadecimal 0x5257CA9A, or dotted-decimal 82.87.202.154. Consequently, the associated 6to4 site prefix is 2002:5257:ca9a::/48. If the stated subnet value and interface identifier are retained, an address in that 6to4 site could be written as 2002:5257:ca9a:1:230:65ff:fe2:9a6. The proposed address does not use the required 2002 6to4 prefix and does not embed the hexadecimal representation of the supplied IPv4 value. RFC 3056 defines this 6to4 address construction, including the use of the IPv4 address immediately after the 2002 prefix. Huawei tunnel implementations follow this standards-based encoding when configuring 6to4 automatic tunnels. See [RFC 3056](#) and [RFC 4219](#).

QUESTION NO: 18

Which of the following options uses an alternate channel to enable communication between IPv6 addresses?

- A. DualStack
- B. 6to4
- C. ISATAP
- D. NAT64

ANSWER: B C

Explanation:

6to4 and ISATAP are IPv6 transition mechanisms that use an IPv4 network as an alternate transport channel for IPv6 traffic. This allows separated IPv6-capable hosts or networks to communicate while native IPv6 connectivity is not available end to end. 6to4 is an automatic tunneling mechanism: it derives an IPv6 prefix from a public IPv4 address and encapsulates IPv6 packets inside IPv4 packets for transit across an IPv4 infrastructure. This supports connectivity between 6to4 sites through an IPv4-based network.

ISATAP (Intra-Site Automatic Tunnel Addressing Protocol) likewise encapsulates IPv6 packets in IPv4, but is designed primarily for communication within an IPv4 intranet. It treats the IPv4 network as a non-broadcast multiple-access link, enabling IPv6 nodes to form tunnel endpoints automatically using IPv4 addressing information. In both cases, the IPv4 packet delivery path is the alternate channel carrying the encapsulated IPv6 packet, which directly matches the question's description. Huawei's IPv6 transition documentation describes IPv6-over-IPv4 tunneling methods, and the IETF specifications define the operational behavior of these mechanisms. See [RFC 3056: Connection of IPv6 Domains via IPv4 Clouds](#) and [RFC 5214: Intra-Site Automatic Tunnel Addressing Protocol](#).

QUESTION NO: 19

Which of the following functions can DHCP snooping provide on a Huawei switch?

- A. It can discard DHCP server messages received on untrusted interfaces.
- B. It can generate DHCP snooping binding entries for valid clients.
- C. It encrypts DHCP messages between the client and server.
- D. It replaces the DHCP server function on the network.

ANSWER: A B

Explanation:

DHCP snooping distinguishes trusted from untrusted interfaces. DHCP server-side messages, such as DHCPOFFER and DHCPACK messages, received on an untrusted interface can be discarded. This prevents an unauthorized DHCP server connected to an access port from assigning incorrect addresses, gateways, or DNS information to users.

DHCP snooping can also build a binding table by recording valid DHCP lease information, including client MAC address, assigned IP address, VLAN, lease time, and interface. Other security functions, such as IP source check or dynamic ARP inspection where supported, can use this trusted binding information to validate user traffic. DHCP message behavior is specified in [RFC 2131](#).

QUESTION NO: 20

There are routers as follows, according to which the compromise conclusion is wrong?

- A. R3 must be the DIS of some Level-1 link
- B. The R3 router is a Level-1-2 router
- C. R3 must be the DIS of some Level-2 link
- D. The system ID of the R3 router is 000300000000

ANSWER: A

Explanation:

R3 must be the DIS of some Level-1 link is the incorrect conclusion. In IS-IS, the Designated Intermediate System role is elected independently on each broadcast network and separately for each routing level. A router's presence in an IS-IS topology, its system ID, or its capability to operate at more than one level does not by itself establish that it is the elected Level-1 DIS on any particular LAN. The Level-1 DIS is determined from the routers participating in Level-1 adjacency on that LAN, primarily using the configured IS-IS interface priority and then the System ID as the election tie-breaker. It is therefore not valid to state that R3 *must* be a Level-1 DIS without the relevant per-interface Level-1 election information. The DIS role also has local-link scope rather than being a router-wide role: the same router can be DIS on one LAN and not be DIS on another. This behavior is defined by the IS-IS LAN election procedures in [RFC 1195](#) and is consistent with Huawei's IS-IS configuration and operational guidance in the [Huawei IP Routing Configuration Guide](#).

QUESTION NO: 21

Configuration under a certain interface: isis timer hello 5 level-2, the following statement is correct?

- A. The hello packets sent at level-1 and Level-2 are both 5s
- B. The hello packets sent at level-1 are sent at 5s interval
- C. The level-2 hello packet send interval is 5s
- D. The CSNP packet sending interval for this interface Level-2 is 5S

ANSWER: C

Explanation:

The statement “*The level-2 hello packet send interval is 5s*” is correct. In Huawei VRP, the interface-view command `isis timer hello 5 level-2` explicitly configures the transmission interval of Level-2 IS-IS Hello (IIH) packets on that interface as 5 seconds. The `level-2` keyword is significant: it scopes the configured timer to the Level-2 IS-IS adjacency process rather than applying the value to all IS-IS levels. IIH packets are used by IS-IS routers to discover neighbors and maintain established adjacencies. Therefore, after this configuration is applied, the device periodically sends Level-2 hello packets every five seconds on the relevant interface, supporting Level-2 neighbor maintenance and liveness detection. The value is an IIH hello interval, not a generic timer for other IS-IS control packets. Huawei’s IS-IS command documentation describes `isis timer hello` as the command used to set the interval for sending IS-IS hello packets, with a `level` keyword available to select the applicable IS-IS level. See Huawei’s [VRP documentation portal](#) and the IS-IS protocol specification, [RFC 1195](#).

QUESTION NO: 22

When deploying BGP/MPLS IPVPN, OSPF's VPN route tag is not present

The extended properties of MP-BGP are passed only locally, only locally, and only meaningful on routers that receive MP-BGP routes and generate OSPF LSAs.

- A. TRUE
- B. FALSE

ANSWER: A

Explanation:

TRUE is correct. In an OSPF-based BGP/MPLS IP VPN, the OSPF route tag is represented by the OSPF Route Tag Extended Community when OSPF information is advertised through MP-BGP. This extended community is non-transitive, so its scope is local rather than end-to-end across subsequent BGP propagation. Its purpose is to preserve the OSPF external route-tag value for the PE that receives the VPN route through MP-BGP and then redistributes or originates that route into an OSPF domain.

This behavior is important because the receiving PE must construct the appropriate OSPF LSA and retain the route-tag information associated with the external OSPF route. The tag therefore has operational significance specifically at the boundary where a VPN route learned by MP-BGP is converted back into OSPF routing information. It is not intended to become a globally propagated VPN route attribute beyond that local redistribution function. RFC 4577 defines the OSPF Route Tag Extended Community and specifies its non-transitive behavior for OSPF support in BGP/MPLS IP VPNs: [RFC 4577](#). Huawei’s MPLS VPN documentation also describes the use of OSPF route attributes during PE-to-CE OSPF route exchange: [Huawei Enterprise Documentation](#).

QUESTION NO: 23

The correct description about the BGP reflector is ?(Multiple choice questions).

- A. IBGP neighbor relationships need to be fully interconnected without a route reflector. The introduction of route reflectors can reduce the need for full interconnection
- B. The route reflector can advertise routes learned from non-clients to all clients
- C. The route reflector can advertise routes learned from one client to other clients and non-clients
- D. The route reflector can advertise routes learned from IBGP neighbors to all clients and non-client

ANSWER: A B C

Explanation:

A route reflector removes the need for every IBGP speaker in an autonomous system to establish a direct IBGP session with every other speaker. Thus, “IBGP neighbor relationships need to be fully interconnected without a route reflector. The introduction of route reflectors can reduce the need for full interconnection” correctly describes the main scaling benefit of route reflection. A reflector applies distinct advertisement rules according to whether the route was received from a client or a non-client. Routes received from a non-client IBGP peer are reflected to the reflector’s clients, so “The route reflector can advertise routes learned from non-clients to all clients” is correct. Routes received from a client are reflected to the reflector’s other clients and to its non-client peers, making “The route reflector can advertise routes learned from one client to other clients and non-clients” correct. These reflection rules allow a hierarchical IBGP design while preserving loop-prevention information through the ORIGINATOR_ID and CLUSTER_LIST attributes. The behavior is standardized in [RFC 4456: BGP Route Reflection](#); the underlying IBGP advertisement model is defined by [RFC 4271: BGP-4](#).

QUESTION NO: 24

What is the security level for untrust zones in the USG Series firewalls? (Radio).

- A. 50
- B. 5
- C. 10
- D. 15

ANSWER: B

Explanation:

5 is the default security level of the *untrust* security zone on Huawei USG Series firewalls. Huawei firewalls use security zones to represent relative trust levels and assign each predefined zone a default numerical level. The untrust zone is intended for networks with the lowest trust level, most commonly the public Internet or an external network. Its level of 5 allows the firewall to use zone security levels as a baseline when evaluating traffic direction and applying the appropriate security policy behavior.

In the standard predefined zone model, the local zone has the highest level, followed by the trust zone and the DMZ zone, while the untrust zone has the lowest level. Administrators can create policies between zones to explicitly permit required services, regardless of the zones' relative security levels. Knowing the default untrust value is important when designing Internet-edge security policies, NAT configurations, and controlled access from internal or DMZ resources to external destinations. Huawei documentation and product support resources covering firewall security-zone configuration are available through the [Huawei Enterprise Support portal](#) and Huawei's [documentation search for USG security-zone levels](#).

QUESTION NO: 25

In an intra-domain MPLSVPN network, when a packet enters the public network and is forwarded, it is encapsulated with two layers of MPLS tags, and the description of the two layers of tags in the following options is incorrect

- A. The outer label of mpls VPN is assigned by the LDP protocol or statically, and the inner label is assigned by the MP-BGP neighbor of the peer
- B. The outer label of MPLS VPN is called a private network label, and the inner label is called a public network label
- C. The outer label is used to correctly send packets to the appropriate VPN in
- D. on the PE device By default, the outer label is ejected before the packet is forwarded to the last hop device

ANSWER: B C

Explanation:

The incorrect descriptions are “The outer label of MPLS VPN is called a private network label, and the inner label is called a public network label” and “The outer label is used to correctly send packets to the appropriate VPN in.” In the standard MPLS

Layer 3 VPN forwarding model, the label stack has distinct transport and service functions. The outer MPLS label is the public-network, or transport, label. It is used by provider-core and provider-edge routers to forward the packet along the label-switched path toward the egress PE. It is therefore not the private-network label. The inner MPLS label is the VPN, or service, label advertised for the VPN route by the egress PE. Once the packet reaches that PE, this inner label identifies the relevant VPN forwarding context/VRF and enables the PE to make the correct customer-side forwarding decision. Thus, associating the outer label with selection of the appropriate VPN reverses the actual functions of the two labels. This two-label behavior is defined by the BGP/MPLS IP VPN architecture: the BGP-distributed VPN label identifies the egress VPN forwarding information, while the transport label carries the packet through the provider backbone. See [RFC 4364](#) and [RFC 3031](#).

QUESTION NO: 26

Router RL and Router R2 run BGP both routers are in the AS65234 The router R2 route is present in the BGP routing table of router RL, but not in the P of router R1 road

- A. With BGP, both routers are in the AS65234 Route cry PD is done by the right big worker
- B. By the table, so what causes the problem to occur? A, synchronization is turned off
- C. The B. BGP neighbor relationship is Down
- D. The route is not optimal.

ANSWER: D

Explanation:

The route is not optimal. A route can be present in the BGP routing table because Router R1 has received and accepted the advertisement from Router R2, while still not being installed in Router R1's IP routing table. Before installation, BGP performs best-path selection among all valid paths for the same destination prefix. Only the selected best BGP path is eligible to be submitted to the global IP routing table. If another path to that prefix has a better BGP path-selection result, the path learned from Router R2 remains visible in BGP but is not used as the active forwarding route. This distinction is fundamental when troubleshooting BGP: the BGP table records candidate paths, whereas the IP routing table contains the route selected for forwarding. The standard BGP decision process and best-path concept are defined in [RFC 4271](#), particularly the route selection rules. Huawei's product documentation portal also provides VRP BGP configuration and route-selection guidance at [Huawei Enterprise Support Documentation](#).

QUESTION NO: 27

There is an existing switch running the RSTP protocol. If the network topology changes, what happens to the Layer 2 forwarding table that is automatically learned by the switch?

- A. All table entries are dropped
- B. Except for the table entry related to the edge port that is not deleted, all other table entries are deleted
- C. Only table entries related to the port that received the TC message are dropped
- D. If the aging time is set to 15 seconds, entries that exceed the aging time are deleted
- E. Table entries are dropped except for table entries related to edge ports and table entries related to ports that receive TC messages

ANSWER: E

Explanation:

Table entries are dropped except for table entries related to edge ports and table entries related to ports that receive TC messages is correct. RSTP uses topology-change notifications to ensure that dynamically learned Layer 2 forwarding database entries do not continue to point traffic toward an old topology after a link or port-state change. When a

switch receives a TC message, it removes dynamically learned MAC-address entries associated with its relevant non-edge ports, allowing those addresses to be relearned promptly through the current active forwarding paths. This prevents frames from being forwarded based on stale MAC-to-port mappings and supports RSTP's rapid convergence behavior.

The received-TC port is retained because it is the direction from which the topology-change notification arrived; retaining entries associated with that interface avoids unnecessarily disrupting reachability in that direction. Entries learned on edge ports are also retained. Edge ports connect to end devices rather than other bridges, so their MAC learning is not normally affected by an RSTP topology change elsewhere in the bridged network. This selective flushing is the Huawei RSTP behavior for handling TC information. See Huawei's [RSTP overview](#) and the [Huawei switching documentation](#).

QUESTION NO: 28

If multiple candidate RPs are configured in a multicast group, which of the following parameters does the RP that elect the group from multiple candidate RPs need to be compared?

- A. The mask length of the group range for the C-RP service that matches the group address that the user joined
- B. C-RP priority
- C. IP address of the C-RP interface
- D. C-RP interface

ANSWER: A B C

Explanation:

In PIM-SM Bootstrap Router (BSR) candidate-RP selection, the elected RP for a multicast group is determined through an ordered comparison of the candidate RPs whose advertised group ranges match that group. First, the group-range mask length is evaluated: the candidate RP advertising the most specific matching group range, meaning the longest mask, takes precedence. This permits different RPs to be selected predictably for more-specific multicast address ranges within a broader configured range.

When eligible candidate RPs have the same matching mask length, their C-RP priorities are compared. Priority provides an administrative way to prefer one RP over another. If both the group-mask length and priority are equal, the IP address advertised by the C-RP is used as the final deterministic tie-breaker; the candidate with the higher address is selected. Therefore, the group-range mask length, C-RP priority, and the C-RP interface IP address are all required comparison parameters. This selection process ensures that all PIM routers derive the same RP mapping from BSR information. The standardized BSR candidate-RP selection rules are specified in [RFC 5059](#); Huawei's multicast implementation follows PIM-SM BSR operational principles described in its [IP Multicast configuration documentation](#).

QUESTION NO: 29

Is the following description of the differences between IGMP name versions correct?

- A. IGMPv1/v2 cannot elect the querier itself, while IGMPv3 can
- B. IGMPv1 does not support specific group queries, while IGMPv2 does
- C. IGMPv1/2/v3 cannot support the SSM model
- D. For members to leave, IGMPv2/v3 can actively leave, while IGMPv1 cannot

ANSWER: B D

Explanation:

IGMPv1 does not support specific group queries, while IGMPv2 does is correct because IGMPv1 defines only the general membership query, which is sent to all hosts and multicast groups on a subnet. IGMPv2 adds the Group-Specific Membership Query, enabling the querier to verify whether receivers remain interested in one particular multicast group. This

makes leave processing more efficient because the router can check only the group affected rather than waiting for the general-query process.

For members to leave, IGMPv2/v3 can actively leave, while IGMPv1 cannot is also correct. IGMPv2 introduces the Leave Group message, allowing a host to explicitly notify multicast routers when it stops receiving a group. IGMPv3 similarly supports explicit membership state changes; when leaving a group, a host reports a state that indicates it no longer wishes to receive the relevant multicast traffic. These mechanisms allow routers to begin group-membership verification promptly, instead of relying solely on membership-report timeout behavior. The protocol details are defined in [RFC 2236](#) for IGMPv2 and [RFC 3376](#) for IGMPv3.

QUESTION NO: 30

You try to trace a remote Internet address on your PC, but the traceroute stops after it reaches a router, and currently, the router should be in the incoming direction of the serial interface

With an ACL "rule5 permit tcp source any destination any", you need to add which of the following ACLs to the router for traceoute to work

- A. Rule 10 permit icmp source any destination any icmp-type ttl-exceeded; rule 15 permit icmp source any destination any icmp-type echo-reply
- B. Rule 10 permit icmp source any destination any icmp-type echo rule 15 permit icmp source any destination any icmp-type net-unreachable
- C. Rule 10 permit tcp source any destination any
- D. Rule 10 permit icmp source any destination any icmp tti-exceeded rule permit icmp source any icmp-type port-unreachableE. Rule 10 permit udp source any destination any rule 15 permit icmp source any destination any icmp-type protocolunreachable

ANSWER: A

Explanation:

Rule 10 permit icmp source any destination any icmp-type ttl-exceeded; rule 15 permit icmp source any destination any icmp-type echo-reply is correct because a PC traceroute based on ICMP Echo packets requires ICMP control messages to return through the router's inbound serial interface. Traceroute sends probes with successively larger TTL values. Each intermediate router that decrements a probe TTL to zero creates an ICMP Time Exceeded message and sends it back to the source. Allowing the `ttl-exceeded` ICMP type therefore lets the PC receive the hop-by-hop responses needed to display the path. When the probe eventually reaches the destination, the destination responds to an ICMP Echo request with an ICMP Echo Reply. Allowing `echo-reply` lets the PC receive that final response and complete the trace. The existing TCP permit rule does not cover these ICMP control packets, so explicit ICMP ACL entries are needed. Microsoft documents that Windows `tracert` uses ICMP Echo Request messages, while ICMP defines both Echo Reply and Time Exceeded message behavior. See [Microsoft tracert documentation](#) and [RFC 792: Internet Control Message Protocol](#).

QUESTION NO: 31

The correct statement about the difference between OSPFv2 and OSPFv3 is ? Multi-select

- A. BOTH OSPFv2 and OSPFv3 support interface authentication. And OSPFv3 authentication is still based on the fields in the ospf message header
- B. OSPFv3 is similar to OSPFv2 in that it uses a multicast address as the OSPF message destination address
- C. The same type of message exists in OSPFv2, OSPFv3: Hello, DD, LSR, LSU, LSAack, Their message format is the same
- D. OSPFv2 has the protocol number 89 in the IPv4 packet header, OSPFv 3 The next number in the IPv6 message header is 89

ANSWER: B D

Explanation:

"OSPFv3 is similar to OSPFv2 in that it uses a multicast address as the OSPF message destination address" is correct. Both protocol versions use link-local-scope multicast to communicate with neighboring routers: the AllSPFRouters and AllDRouters groups. For IPv4-based OSPFv2 these are 224.0.0.5 and 224.0.0.6; for IPv6-based OSPFv3 they are FF02::5 and FF02::6. This preserves the operational model of sending Hello packets and other applicable control traffic to local OSPF participants rather than requiring separate unicast transmissions to every neighbor.

"OSPFv2 has the protocol number 89 in the IPv4 packet header, OSPFv3 The next number in the IPv6 message header is 89" is also correct. OSPF is assigned IP protocol number 89. In an IPv4 packet, this value is carried by the Protocol field. In an IPv6 packet, it is carried by the Next Header field, subject to normal IPv6 extension-header processing. OSPFv3 changed its addressing and packet-header design for IPv6, but it retained the assigned protocol value. The OSPFv3 specification documents the IPv6 multicast groups and packet encapsulation, while the IANA protocol-number registry lists OSPF as 89. See [RFC 5340](#) and [IANA Assigned Internet Protocol Numbers](#).

QUESTION NO: 32

By default, routers running IS-IS in broadcast networks have a period of seconds for DIS to send CSNP packets?

- A. 33
- B. 40
- C. 30
- D. 10

ANSWER: D**Explanation:**

10 is correct. On an IS-IS broadcast multi-access network, the Designated Intermediate System (DIS) is responsible for periodically transmitting Complete Sequence Number PDUs (CSNPs). In Huawei VRP, the default CSNP interval on a broadcast interface is 10 seconds. These periodic CSNPs provide the database-synchronization mechanism for the LAN: they summarize the link-state database known by the DIS, allowing other IS-IS routers on the segment to identify missing or obsolete Link State PDUs and request or reconcile the required information. This behavior is particularly important after adjacency establishment and while maintaining consistent link-state databases among all routers attached to the broadcast segment. The interval can be tuned with the relevant IS-IS CSNP interval configuration command when a design requires a different synchronization frequency, but without explicit configuration Huawei uses the 10-second default. IS-IS CSNP operation on broadcast networks is defined in the ISO IS-IS extensions for IP, including the DIS role and periodic database summaries; see [RFC 1195](#). Huawei product documentation and VRP command references are available through the [Huawei Enterprise Documentation portal](#).

QUESTION NO: 33

What is the field that the IPv6 header has similar to the IPv4 header "Type of Service" field?

- A. Version
- B. Flow Label
- C. Next Header
- D. Traffic Class

ANSWER: D**Explanation:**

Traffic Class is the IPv6 base-header field corresponding to the IPv4 Type of Service field. It is an 8-bit field used to identify the desired handling of a packet as it crosses a network. In contemporary IP terminology, these bits carry Differentiated

Services Code Point (DSCP) information for differentiated-services quality-of-service classification and Explicit Congestion Notification (ECN) information for congestion signaling. This enables routers and other network devices to classify traffic and apply forwarding behavior, such as prioritizing latency-sensitive voice or video traffic according to an established QoS policy.

IPv6 retained this traffic-handling capability while simplifying the overall base header. The IPv6 Traffic Class field has the same 8-bit size and intended DSCP/ECN usage as the IPv4 differentiated-services interpretation of the former Type of Service octet. This direct functional correspondence is specified in the IPv6 header format defined by the IETF. Huawei QoS deployments likewise use IP precedence or DSCP-based classification to associate packets with service classes and forwarding behavior. See [RFC 8200, IPv6 Header Format](#) and [RFC 2474, Definition of the Differentiated Services Field](#).

QUESTION NO: 34

The following description is correct regarding route control

- A. IBGP routes can be introduced using the import-route BGP command.
- B. AS-path-filter is used to match AS paths and affects route control.
- C. OSPF, IS-IS, and BGP can advertise default routes.
- D. In a route-policy node, if-match community-filter 1 and if-match as-path-filter 1 have an AND relationship.

ANSWER: B C D

Explanation:

AS-path filtering is a valid BGP route-control mechanism. In Huawei VRP, an AS-path filter defines matching rules against the BGP AS_Path attribute, and that filter can be referenced by a route policy to match routes and control whether they are accepted, advertised, preferred, or assigned attributes. This enables policy decisions based on the autonomous systems a route has traversed, which is central to interdomain path control.

OSPF, IS-IS, and BGP each support mechanisms for originating or advertising a default route. In operational practice, this lets a routing domain direct traffic for destinations lacking a more-specific route toward an intended exit or upstream path. The precise command and prerequisites vary by protocol and VRP feature context, but default-route advertisement is supported by all three protocols.

Within a Huawei route-policy node, different `if-match` clauses are evaluated with an AND relationship. Consequently, using both a community-filter match and an AS-path-filter match requires a route to satisfy both conditions before that node matches. This supports precise, layered policy control using multiple independent BGP attributes. See the [Huawei Enterprise documentation portal](#) and the BGP AS-path definition in [RFC 4271](#).

QUESTION NO: 35

The following description of Link ID, Data, Type, and Metric in Router-LSA is correct : (Multiple select).

- A. Metric describes the overhead of this connection
- B. The Link ID represents the local identity of this connection, and the different connection types The Link ID represents different meanings
- C. Data is used to describe additional information for this connection, and different connection types describe different information
- D. Type indicates the type of connection

ANSWER: A B C D

Explanation:

All four statements correctly describe the fields of an OSPF Router-LSA link description. A Router-LSA, originated by each router for its area, contains one or more link entries that advertise the router's connections to the OSPF topology. The metric

field is the OSPF cost associated with reaching the advertised link or network through that connection; it is therefore the connection overhead used in shortest-path calculation. The link ID is an identifier whose precise value is determined by the link type. For example, it can identify a neighboring router, the designated router for a transit network, or a stub network destination. Link data supplements that identifier with type-specific information, such as the advertising router interface address or a network mask. The type field explicitly defines the nature of the advertised connection, allowing OSPF to interpret the link ID and link data appropriately. These fields work together so that routers in the area can construct a consistent link-state database and run the SPF algorithm using the advertised costs. The Router-LSA link-description format and the meanings of these fields are specified in [RFC 2328, section 12.4.1](#).

QUESTION NO: 36

Regarding the WRR (Weight Round Robin) description of the error is ?(Multiple select).

- A. WRR avoids congestion in the network
- B. WRR guarantees that various queues are allocated to a certain amount of bandwidth
- C. WRR is a congestion management technique
- D. WRR can no longer be used
- E. WRR on GE interfaces ensures that critical services are prioritized

ANSWER: A D E

Explanation:

The erroneous descriptions are “WRR avoids congestion in the network,” “WRR can no longer be used,” and “WRR on GE interfaces ensures that critical services are prioritized.” WRR is an egress-queue scheduling mechanism used when traffic contends for limited outbound interface bandwidth. It manages that contention by servicing queues according to configured weights; it does not prevent congestion from arising elsewhere in a network or remove the need for capacity planning, traffic policing, and shaping. WRR remains a supported QoS scheduling method in Huawei VRP-based platforms, so a blanket assertion that it can no longer be used is inaccurate. In addition, WRR distributes service among queues proportionally to their configured weights and does not provide strict, absolute precedence for a critical queue. Where critical, delay-sensitive traffic requires preferential forwarding, Huawei QoS designs use strict-priority scheduling, commonly in a combined PQ+WRR scheduling model, together with appropriate rate controls. Huawei’s enterprise support portal provides VRP QoS configuration and scheduling documentation at [Huawei Enterprise Documentation](#). The general behavior and operational considerations of weighted round-robin scheduling are also described in [RFC 4594](#).

QUESTION NO: 37

What are the problems with manually creating static VXLAN tunnels in campus networks? Multi-select

- A. Although the static VXLAN tunnel mode can support distributed gateway application scenarios, the configuration workload is large and the configuration adjustment is complex
- B. N devices to establish static VXLAN tunnels, you need to manually configure up to $N(N-1)/2$ tunnels, the amount of configuration is large
- C. VTEP can only learn remote MAC addresses in a way that relies on data flooding
- D. Static VLAN tunnel also uses related protocols on the control plane, which will cause equipment resource loss

ANSWER: A B C

Explanation:

“Although the static VXLAN tunnel mode can support distributed gateway application scenarios, the configuration workload is large and the configuration adjustment is complex” is correct because a manually built VXLAN overlay requires tunnel peers and related forwarding parameters to be maintained device by device. As endpoints, VTEPs, or tenant segments change,

those settings must be updated consistently throughout the affected fabric, which makes expansion and operational changes cumbersome.

“N devices to establish static VXLAN tunnels, you need to manually configure up to $N(N-1)/2$ tunnels, the amount of configuration is large” correctly describes the full-mesh scalability issue. Each pair of VTEPs requires a tunnel relationship, so the number of required pairwise tunnels grows quadratically as more VTEPs are introduced.

“VTEP can only learn remote MAC addresses in a way that relies on data flooding” is also correct for static VXLAN operation without a control-plane MAC distribution mechanism. Unknown-destination traffic and broadcast or multicast traffic must be replicated through the configured VXLAN peers, and remote MAC learning occurs from the encapsulated data-plane traffic. EVPN is designed to distribute MAC/IP reachability through BGP and reduce this dependence on flooding. See Huawei’s [VXLAN Configuration Guide](#) and the [VXLAN EVPN configuration documentation](#).

QUESTION NO: 38

BGP authentication can prevent illegal routers from establishing neighbors with BGP routers, BGP authentication can be divided into MD5 authentication and Keychain authentication, please ask which BGP messages will be carried below BGPKeychain authentication information?(Multiple select).

- A. Update
- B. Notification
- C. Keepalive
- D. Open

ANSWER: A B C D

Explanation:

BGP Keychain authentication protects the TCP connection used by a BGP peer relationship. BGP protocol messages are transported over that authenticated TCP session, so the authentication information is applied to TCP segments carrying UPDATE, NOTIFICATION, KEEPALIVE, and OPEN messages. Therefore, all four listed BGP message types are covered when BGP Keychain authentication is configured. The OPEN message is especially important because it is exchanged while the BGP session is being established; protecting it ensures that peer negotiation occurs over an authenticated connection. After the peer relationship is established, KEEPALIVE messages maintain the session, UPDATE messages advertise or withdraw routing information, and NOTIFICATION messages report errors and close a session when necessary. Authentication must continue to protect each of these message types so that a device can validate that received BGP traffic belongs to the expected peer and has not been injected or altered in transit. This behavior aligns with BGP’s use of TCP as its transport protocol and with TCP authentication mechanisms, which protect TCP segments rather than only a selected BGP message type. See the BGP message definitions in [RFC 4271](#) and the TCP Authentication Option specification in [RFC 5925](#).

QUESTION NO: 39

The following statement about IPSec VPN is correct ?

- A. Two routers establish an IPSec VPN in brutal mode, so that the two devices can interact with a minimum of 4 messages to establish a tunnel.
- B. Brutal mode can support NAT traversal, while main mode does not support NAT traversal.
- C. The two routers can establish OSPF neighbor relations and exchange intranet routes through the IPSec VPN tunnel
- D. The two routers establish IPSecVPN in main mode, and from the 5th packet (included), the payload's data will be encrypted.

ANSWER: D

Explanation:

The statement “The two routers establish IPsecVPN in main mode, and from the 5th packet (included), the payload's data will be encrypted.” is correct for IKE version 1 Main Mode negotiation. Main Mode uses six IKE messages to establish the ISAKMP/IKE security association. The first two messages negotiate the IKE security-association proposal, and the third and fourth messages exchange Diffie-Hellman public values and nonces. At that point, both peers have enough shared keying material to derive the encryption and integrity keys protecting subsequent IKE exchanges. Therefore, the fifth and sixth Main Mode messages, which carry the peer identity and authentication information, have encrypted payloads. This design protects the identities of the VPN peers during the initial negotiation, an important Main Mode characteristic. The wording concerns encryption of IKE message payloads beginning with the fifth message, rather than the later encrypted user traffic carried after IPsec SAs are established. The IKEv1 Main Mode exchange and the encryption of the final identity/authentication messages are specified in [RFC 2409, section 5](#). Huawei IPsec implementations that use IKEv1 Main Mode follow this six-message negotiation behavior; see Huawei's [IPsec VPN documentation](#) for implementation guidance.

QUESTION NO: 40

Is the following description of the BSR/RP mechanism correct?

- A.** A PIM-SM domain can have more than one C-BSR, but only one BSR is elected. The elected BSR receives C-RP packets and collects C-RP information.
- B.** The BSR advertises BSR and C-RP information to all routers in the PIM-SM domain by flooding Bootstrap packets.
- C.** A C-BSR can also collect C-RP information by receiving C-RP packets.

ANSWER: A B C**Explanation:**

In a PIM-SM domain, multiple routers may be configured as Candidate Bootstrap Routers (C-BSRs), providing redundancy for the bootstrap function. These candidates participate in the BSR election, but the election produces one active Bootstrap Router (BSR). The active BSR receives Candidate-RP-Advertisement messages from Candidate RPs, builds the group-to-RP mapping information, and distributes that information throughout the PIM-SM domain. This behavior is defined by the PIM Bootstrap Router mechanism in [RFC 5059](#).

The selected BSR sends Bootstrap messages hop by hop across the PIM-SM domain. PIM routers process and forward these Bootstrap messages, which effectively floods the BSR address and RP-set information to all PIM routers. Routers then use the learned RP-set to select an RP for relevant multicast groups. A C-BSR is capable of receiving Candidate-RP-Advertisement messages and collecting Candidate-RP information; when it becomes the elected BSR, that collected information is used to generate and advertise the RP-set. The broader PIM-SM control-plane behavior, including the RP role in source registration and shared-tree construction, is specified in [RFC 7761](#).

QUESTION NO: 41

What messages can Huawei equipment suppress traffic for?

- A.** Multicast
- B.** Known unicast
- C.** Broadcast
- D.** Unknown unicast

ANSWER: A C D**Explanation:**

Huawei switch traffic suppression, often configured as storm control, applies to the flooding-prone Layer 2 traffic categories: multicast traffic, broadcast traffic, and unknown-unicast traffic. These frames can be replicated or flooded across a VLAN

and, if their rate becomes excessive, consume link bandwidth and device forwarding resources. Suppression allows an administrator to configure a rate threshold for each applicable traffic type on an interface. When measured traffic exceeds the configured threshold, the device discards excess frames of that type, helping contain broadcast storms, multicast flooding, and unknown-unicast flooding while preserving normal network availability. The behavior is particularly useful at access-facing interfaces, where loops, faulty endpoints, or abnormal application behavior can rapidly generate large volumes of flooded frames. Huawei documentation describes traffic suppression as controlling broadcast, multicast, and unknown-unicast packets according to configured thresholds; the relevant feature and command documentation can be located through the [Huawei Enterprise Support traffic-suppression search](#) and the [Huawei Enterprise Support documentation portal](#). Therefore, the applicable message categories are multicast, broadcast, and unknown unicast.

QUESTION NO: 42

As shown in the following figure, a part of the LSDB of a router is shown in the first one in the following figure. A new LSP has been received, the second one in the figure below, and the following statement is incorrect (single choice).

- A. this router will put the newly received LSP into the LSDB
- B. In the case of a broadcast network, DIS will include summary information for this LSP in the next CSNP message.
- C. This router ignores LSPs received from neighbors
- D. If it is a point-to-point network, the router will send PSNP

ANSWER: C

Explanation:

The incorrect statement is *"This router ignores LSPs received from neighbors."* IS-IS routers use the LSP identifier, sequence number, checksum, and remaining lifetime to compare a received LSP with the corresponding LSDB entry. When the received LSP is newer than the stored copy—as represented by the new LSP in the figure—the router installs that newer version in its LSDB rather than discarding it. This LSDB update is fundamental to IS-IS link-state synchronization: the router must retain the most current information so that its shortest-path calculation uses the current topology and reachability data. After accepting the LSP, the normal acknowledgment and database-synchronization behavior depends on the network type. On a point-to-point circuit, a PSNP can acknowledge receipt; on a broadcast circuit, the DIS periodically advertises database summaries in CSNPs. These procedures support reliable flooding, but the key comparison result is that a newer received LSP must replace the older LSDB copy. The IS-IS protocol's LSP comparison and database synchronization behavior is specified in [RFC 1195](#); point-to-point acknowledgement behavior is further described in [RFC 5303](#).

QUESTION NO: 43

What information does the BGP Open message carry as follows?

- A. BGP Router ID
- B. Route attributes
- C. Local Autonomous System (AS) number
- D. Hold time

ANSWER: A C D

Explanation:

The BGP Open message establishes the essential session parameters that two peers must exchange before they can create a BGP adjacency. "Local Autonomous System (AS) number" is carried in the Open message as the sender's autonomous-system value. In the original protocol format, this is the two-octet *My Autonomous System* field; support for four-octet autonomous-system values is negotiated using an Open-message capability.

“Hold time” is also a mandatory Open-message field. Each peer advertises its proposed hold-time value, allowing both sides to derive the negotiated timer used to detect loss of the BGP connection when expected keepalive or update traffic is not received.

“BGP Router ID” is conveyed as the BGP Identifier field in the Open message. This four-octet value identifies the BGP speaker and must be unique within the autonomous system. Together, these values let peers validate compatibility and establish the operational parameters for the BGP session before exchanging routing information. The message format and field definitions are specified in [RFC 4271, section 4.2](#); four-octet autonomous-system capability negotiation is defined by [RFC 6793](#).

QUESTION NO: 44

Which of the following VPN technologies may have two IP message headers in one message?

- A. GOES
- B. SSL VPN
- C. IPsec VPN
- D. L2TP

Answer: ACD

Explanation:

- A. SSL VPN
- B. IPsec VPN
- C. L2TP
- D. GRE

ANSWER: B C D

Explanation:

GRE, IPsec VPN, and L2TP can produce a packet containing an inner and an outer IP header because they encapsulate an original packet for transport across a tunnel. GRE wraps the original Layer 3 packet with a GRE header and a new delivery IP header. The original IP header remains intact as the inner header, while the new header routes the encapsulated packet between tunnel endpoints. Huawei describes GRE as an encapsulation protocol that can carry packets over an IP network; see [Huawei: Understanding GRE](#).

IPsec VPN in tunnel mode protects an entire original IP packet and adds a new outer IP header for routing between the IPsec peers. The protected original packet, including its original IP header, is carried inside the IPsec tunnel. L2TP similarly encapsulates PPP frames in L2TP and UDP for transmission over an IP network. When the PPP payload is an IP packet, its original IP header is retained and an outer IP header is used to deliver the L2TP/UDP packet to the tunnel peer. This encapsulation model is specified in [RFC 2661, Layer Two Tunneling Protocol](#).

QUESTION NO: 45

The following description of MED, correct:

- A. MED default value is 100
- B. MED can only be passed within this AS
- C. MED is an optional non-transitional attribute
- D. MED The default value is 0

ANSWER: C D

Explanation:

“MED is an optional non-transitional attribute” is correct because the Multi-Exit Discriminator is defined by BGP as an optional, non-transitive path attribute. It provides a metric to a neighboring autonomous system to indicate the preferred entry point when multiple inter-AS links exist. As a non-transitive attribute, MED received from an external peer is not propagated onward to another external autonomous system. The BGP base specification defines MED in its path-attribute section and describes its use for selecting among multiple entry points: [RFC 4271, section 5.1.4](#). “MED The default value is 0” is also correct for Huawei VRP BGP behavior. When no MED is explicitly configured or derived by an applicable routing-policy or route-import process, the MED associated with the route uses a value of 0. A lower MED is preferred when BGP compares otherwise eligible routes received from the same neighboring AS, subject to the configured comparison behavior. The general BGP best-path treatment of MED is specified in [RFC 4271, section 9.1.2.2](#).

QUESTION NO: 46

When the BFD detection intervals on both ends are 30ms and 40ms, are the following options described correctly? An A. BFD session can be established, which takes 40ms after negotiation

A B. BFD session can be established, which takes 30ms after negotiation

C. BFD sessions can be established, each occurring at its own interval

D. BFD sessions cannot be established

Answer: C

Explanation:

A. BFD session can be established, which takes 40ms after negotiation

ANSWER:

Explanation:

BFD sessions can be established, each occurring at its own interval is correct. BFD does not require both peers to have an identical configured transmission or detection interval before a session can come up. Instead, the protocol exchanges timing parameters in BFD control packets and determines operation independently in each direction. Each device advertises its desired minimum transmit interval and its required minimum receive interval. The effective sending interval for a direction is derived from the local desired transmit value and the remote peer's receive requirement, so one direction can operate with a different effective interval from the reverse direction. Therefore, endpoints configured with 30 ms and 40 ms values can form a BFD session; there is no requirement to negotiate one common 30 ms or 40 ms interval for both directions. This directional behavior lets BFD accommodate different interface capabilities, platform defaults, and failure-detection requirements while retaining rapid liveness monitoring. The BFD base specification defines these independently advertised timing values and their negotiation behavior in the control-packet and state-machine procedures. Huawei VRP BFD configuration follows this standards-based model through separately configurable transmit and receive timing parameters. See [RFC 5880: Bidirectional Forwarding Detection](#) and [Huawei Enterprise Documentation](#).

QUESTION NO: 47

The following description of the LDPLSP establishment process is correct?

A. By default, LSRs are for the same FEC, and the received tag mappings can only come from the optimal next hop, not from the non-optimal next hop.

B. When a network topology change causes the next hop neighbor to change, the free label keep method is used. LSR can quickly rebuild the LSP by directly using the labels sent by the original non-optimal next-hop neighbor. Liberal needs more memory and tag space.

C. In the DoD mode of label publishing, for a specific FEC, the LSR does not need to obtain a label request message from upstream to distribute the label.

D. The process of establishing a LSP is actually to bind the FEC to the tag and advertise this binding to the LSP upstream LSR process.

ANSWER: B D

Explanation:

When a network topology change causes the next hop neighbor to change, the free label keep method is used. LSR can quickly rebuild the LSP by directly using the labels sent by the original non-optimal next-hop neighbor. Liberal needs more memory and tag space. This accurately describes liberal label retention: an LSR retains label mappings received from peers even when they are not currently the selected next hop for the FEC. If routing later selects one of those peers, the already retained mapping can be used immediately, reducing the time required to restore forwarding. The trade-off is that retaining these additional mappings consumes more label and memory resources. The process of establishing a LSP is actually to bind the FEC to the tag and advertise this binding to the LSP upstream LSR process. This is also correct because LDP distributes FEC-to-label bindings downstream-to-upstream: a downstream LSR assigns a label for a FEC and advertises the mapping to its upstream peer, allowing each hop to build the label-switched forwarding state toward the FEC. These procedures are defined in the LDP specification, including label distribution and retention behavior. See [RFC 5036: LDP Specification](#) and [RFC 5036, Label Retention Mode](#).

QUESTION NO: 48

Which of the following configurations is the correct OSPFv3 route aggregation configuration:

- A. [Huawei] ospfv3 1
[Huawei-ospfv3-1] area 1
[Huawei-ospfv3-1-area-0001] abr-summary fc00:0:0:: 48 cost 400
- B. [Huawei] interface gigabitethernet 1/0/0
[Huawei-GigabitEthernet1/0/0] asbr-summary fc00:0:0:: 48 cost 20
- C. [Huawei] ospfv3 1
[Huawei-ospfv3-1] asbr-summary fc00:0:0:: 48 cost 20 tag 100
- D. [Huawei] ospfv3 1
[Huawei-ospfv3-1] abr-summary fc00:0:0:: 48 cost 400
- E. [Huawei] ospfv3 1
[Huawei-ospfv3-1] area 1
[Huawei-ospfv3-1-area-0001] asbr-summary fc00:0:0:: 48 cost 20 tag 100

ANSWER: A C

Explanation:

The valid OSPFv3 route-aggregation configurations are `abr-summary fc00:0:0:: 48 cost 400` under an OSPFv3 area view and `asbr-summary fc00:0:0:: 48 cost 20 tag 100` under the OSPFv3 process view. The former performs Area Border Router summarization: an ABR advertises one aggregated inter-area IPv6 prefix for routes learned within the specified area. Therefore, the command must be entered after selecting both the OSPFv3 process and the relevant area. The configured cost is advertised with the generated summary route.

The latter performs Autonomous System Boundary Router summarization for redistributed external IPv6 routes. It is configured in the OSPFv3 process view because external-route summarization is an ASBR function rather than an area-specific ABR function. In addition to assigning the summary cost, the configuration can attach an external route tag, here 100, to the summary advertisement. In both commands, the IPv6 prefix and prefix length must be separate arguments, expressed as `fc00:0:0:: 48`. OSPFv3 carries IPv6 routing information and supports the distinct ABR and ASBR summarization roles described by the protocol. See the [OSPF for IPv6 specification \(RFC 5340\)](#) and Huawei's [enterprise product documentation portal](#) for platform-specific command references.

QUESTION NO: 49

Regarding route introduction, the following description is correctly ?(Multiple choice questions).

- A. OSPF introduces two types of external routes, with the first type of external route taking precedence over the second type of external route.
- B. Use the default-route imported command in BGP only to introduce default routes that already exist in the local IP routing table.
- C. When ISIS introduces an external route, if you do not specify the route type that is introduced, the Level-1 route is introduced by default on the L1/2 router
- D. When ISIS introduces an external route, if you do not specify the type of route that is introduced, the Level-1 route is introduced by default on the L1 router .

ANSWER: A B D

Explanation:

“OSPF introduces two types of external routes, with the first type of external route taking precedence over the second type of external route” is correct because OSPF uses Type 1 and Type 2 external metrics. For routes redistributed by multiple ASBRs, a Type 1 external route is preferred to a Type 2 external route; Type 1 incorporates the internal cost to reach the ASBR into the overall path cost. This behavior is specified in the OSPF standard.

“Use the default-route imported command in BGP only to introduce default routes that already exist in the local IP routing table” is also correct in Huawei VRP. This command imports and advertises an existing active default route, rather than unconditionally originating one. It is therefore appropriate when BGP advertisement must depend on the local routing table having a valid default route.

“When ISIS introduces an external route, if you do not specify the type of route that is introduced, the Level-1 route is introduced by default on the L1 router” is correct. Huawei IS-IS imports routes into the level supported by the router when no level is explicitly supplied; on a Level-1 router, the default import level is Level-1. See [RFC 2328, OSPF Version 2](#) and [Huawei Enterprise Documentation](#).

QUESTION NO: 50

Two routers, connected together through the serial port, but can not ping between each other , now check the port status as follows: Can the following information determine what the cause is?

```
[R1]display interface serial 0/0/0
Serial0/0/0 current state : UP
Line protocol current state : DOWN
Description:
Route Port, The Maximum Transmit Unit is 1500, Hold timer is 10(sec)
Internet Address is 22.22.22.2/24
Link Layer protocol is nonstandard HDLC
Last physical up time : 2022-08-01 16:24:00 UTC+08:00
Last physical down time : 2022-08-01 16:24:00 UTC+08:00

[R2]display interface Serial 0/0/0
Serial0/0/0 current state : UP
Line protocol current state : DOWN
Description:
Route Port, The Maximum Transmit Unit is 1500, Hold timer is 10(sec)
Internet Address is 22.22.22.2/24
Link Layer protocol is PPP
LCP stopped
```

- A. Insufficient link bandwidth
- B. The IP addresses on both ends are not on the same network segment
- C. The link layer protocols on both ends are inconsistent
- D. Subnet mask does not match the

Answer: C

Explanation:

- A. The IP addresses on both ends are not on the same network segment
- B. The link layer protocols on both ends are inconsistent
- C. Subnet mask does not match the

ANSWER: B**Explanation:**

The link layer protocols on both ends are inconsistent is correct. For a serial connection, the physical interface can be operational while the line protocol remains unavailable if the peers do not use compatible Layer 2 encapsulation or negotiation settings. A common example is one serial interface using PPP while its peer uses HDLC, or the two ends having incompatible PPP negotiation parameters. In this state, the routers have an electrical and physical serial connection, but cannot form the Layer 2 relationship required to carry IP packets. Consequently, an ICMP ping cannot succeed regardless of the configured IP addresses. The relevant troubleshooting approach is to inspect the serial-interface status and its configured link protocol on both devices, then configure matching encapsulation and any associated parameters at each end. Huawei VRP interface-status output distinguishes the physical state from the protocol state, which is important when isolating this kind of serial-link fault. Huawei documentation and product support resources are available through the [Huawei Enterprise Support portal](#) and the [Huawei documentation portal](#).

QUESTION NO: 51

When DLDP detects the presence of a one-way link in the fiber, the default actions of Huawei equipment include: the DLDP state machine migrates to the Disable state, outputs Trap information, and automatically sets the interface to a blocked state.

- A. True
- B. False

ANSWER: A**Explanation:**

True is correct. Huawei Device Link Detection Protocol (DLDP) is designed to detect unidirectional Ethernet links, a condition that can otherwise allow traffic in one direction while control traffic or return traffic fails. When DLDP determines that an interface is connected through a one-way fiber link, its state machine transitions to the Disable state. The device also generates a trap notification so that network-management systems and administrators can be alerted to the fault condition. As the default protection action, DLDP blocks the affected interface, preventing it from forwarding normal service traffic and therefore helping avoid forwarding loops, MAC-address instability, and traffic black holes caused by asymmetric connectivity. The blocked status is specifically a DLDP protection response; it enables the device to isolate the faulty path while retaining diagnostic visibility into the detected unidirectional-link event. Huawei documents DLDP as a Layer 2 link-fault detection and protection mechanism, including its interface state behavior and alarm reporting. See the Huawei Enterprise Support documentation search for [DLDP](#) and Huawei's [product documentation portal](#) for the applicable switch software release.

QUESTION NO: 52

What is correct below about LDP synchronization in ISIS?

- A. LDP synchronization can be turned on globally.
- B. LDP establishes neighbors faster than ISIS, so LDP synchronization is turned off by default
- C. LDP synchronization avoids traffic loss when the active-standby link is switched.
- D. LDP synchronization can be deployed on the Loopback interface.

ANSWER: C

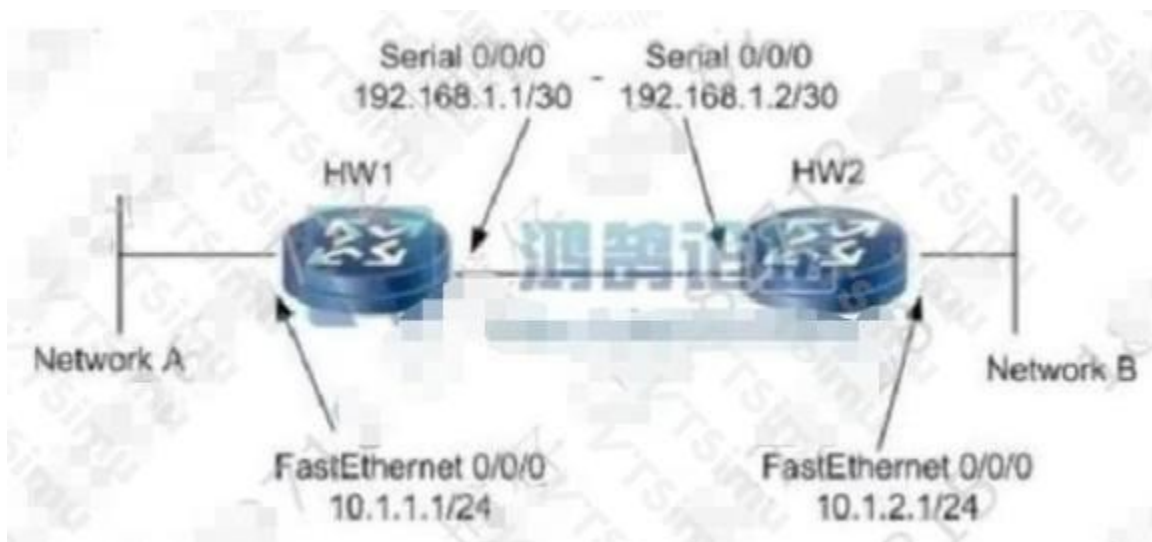
Explanation:

LDP synchronization avoids traffic loss when the active-standby link is switched. IS-IS can establish its adjacency and advertise a newly available link before LDP has completed label-distribution negotiation on that same link. If IS-IS selects the link immediately, MPLS packets may be sent toward a next hop for which the required LDP label forwarding state is not yet available. This creates a transient MPLS forwarding black hole, particularly when traffic moves to a backup path after a failure or switchover.

With LDP-IGP synchronization, IS-IS keeps the relevant link unattractive for transit traffic, typically by advertising it with a maximum metric, until LDP is operational and label bindings can be used for forwarding. Once the LDP session is ready, the normal IS-IS metric is restored and the path can safely carry labeled traffic. This coordinated readiness is exactly why the feature protects MPLS traffic during active-standby path transitions. The behavior is standardized by the LDP IGP Synchronization procedures in [RFC 5443](#); Huawei product documentation is available through the [Huawei Enterprise Support documentation portal](#).

QUESTION NO: 53

The switch can suppress traffic by suppressing traffic and extremely other failures can affect operations, which of the following traffic suppression configurations is the wrong configuration



- A. [Quidway] wan 10 [Quidway-wan10] multicast-suppression 1000 [Quidway-van10] quit
- B. [Quidway] interface gigabitethernet 0/0/1
[Quidway-GigabitEthernet0/0/1] multicast-suppression 80 [Quidway-GigabitEthernet0/0/1] quit
- C. [Quidway] interface gigabitethernet 0/0/1
[Quidway-GigabitEthernet0/0/1] broadcast-suppression 80 [Quidway-GigabitEthernet0/0/1] quit
- D. [Quidway] icmp rate-limit enable [Quidway] icmp rate-limit total threshold 20
- E. [Quidway] wan 10 [Quidway-wan10] broadcast-suppression 1000 [Quidway-van10] quit

ANSWER: A

Explanation:

[Quidway] wan 10 [Quidway-wan10] multicast-suppression 1000 [Quidway-van10] quit is the incorrect configuration because multicast traffic suppression is not configured in VLAN view in this VRP traffic-suppression context. Multicast suppression is a per-interface traffic-control function: the device must apply the configured threshold to the ingress traffic received on the relevant physical Ethernet interface or Eth-Trunk. A VLAN is a Layer 2 broadcast domain and does not itself provide the applicable interface-level context for this multicast-suppression command. Consequently, entering the VLAN configuration context and attempting to configure multicast suppression there is not a valid way to deploy the required

control. In a real network, the engineer should identify the port or aggregated link receiving the excessive multicast traffic, enter that interface view, and configure the permitted multicast threshold according to the platform's supported unit and range. This ensures that the device can measure and discard excess multicast packets where they enter the switching fabric, protecting forwarding and control-plane resources. Huawei command syntax and supported configuration views vary by VRP release and switch series, so the applicable product command reference should always be checked before deployment. See the official [Huawei Enterprise documentation portal](#) and [Huawei Enterprise Support](#).

QUESTION NO: 54

Is the following description of ospf virtual links correct?

- A. A virtual connection can be established in any area, after which it itself belongs to the region
- B. The virtual connection ip is used as the link address
- C. Virtual links can be used to solve the problem of zone 0 being segmented
- D. The cost of a virtual connection is zero and is the optimal link

ANSWER: C

Explanation:

Virtual links can be used to solve the problem of zone 0 being segmented. In OSPF, Area 0 is the backbone and must provide logical connectivity between all non-backbone areas. A virtual link is a logical point-to-point OSPF adjacency configured between two Area Border Routers through a shared non-backbone transit area. It extends backbone connectivity across that transit area, allowing backbone information to be carried when the physical topology has split Area 0 into separate parts. This is a remediation mechanism for a discontinuous backbone, such as after a network merger or where a direct backbone connection is temporarily unavailable. The virtual link is treated operationally as a backbone connection, while the underlying path across the transit area must remain reachable. OSPF derives the virtual-link cost from the intra-area path through the transit area, so it is not inherently a zero-cost or automatically preferred path. The OSPF specification defines virtual links as part of the mechanism for maintaining backbone connectivity and describes their use across a transit area. See [RFC 2328, Section 15](#) and Huawei's [OSPF configuration documentation](#) for implementation context.

QUESTION NO: 55

As shown in the figure, all routers declare the loopback address in OSPF where the R storage does not return address is 1022 [2 declares 10033[at the Ael.R3 ring E port[92101102A External routing in order to win light R win L away from the negative cut.



The noodle command is as follows: (multi-select) [Blpppfixt permlt 1001.1.024 [R1] ospf1

- A. router R1 of the IP routing table, there is no external route to 1001.1.0/24 of
- B. router R1 In the IP route table, there is an external route to 1001.1.0/24
- C. In the IP routing table of Router R1, there is no loopback route to R2
- D. In the IP routing table of Router R1, there is no loopback route to R3

ANSWER: B C D

Explanation:

The configured IP prefix list permits only the prefix 100.1.1.0/24. When it is applied as an OSPF import filtering policy on R1, OSPF installs the permitted external prefix in R1's routing table. Therefore, the statement that R1 has an external route to 100.1.1.0/24 is correct. The route is learned as an OSPF external route because it originates outside the OSPF domain and is advertised by the ASBR through OSPF external LSAs.

The loopback prefixes advertised by R2 and R3 do not match the sole permitted prefix in the configured IP prefix list. As a result, those loopback routes are excluded by the import filter and are not installed in R1's IP routing table. Thus, the statements that R1 has no loopback route to R2 and no loopback route to R3 are also correct. This illustrates that prefix-list matching is exact with respect to the configured prefix length unless a length range is explicitly specified. OSPF route filtering should be designed carefully because it changes which otherwise valid OSPF routes are eligible for installation in the local routing table. See Huawei's [Enterprise Support documentation portal](#) and the OSPF external-route specification in [RFC 2328](#).

QUESTION NO: 56

As shown in the figure, in the context of IPv4 and IPv6, the SEL field in the NET address of isis always takes a value of 00 ?



- A. Right
- B. Wrong

ANSWER: A

Explanation:

Right is correct. An IS-IS Network Entity Title (NET) identifies the IS-IS router itself and is composed of the area address, system ID, and a one-byte NSEL/SEL field. When an NSAP-format address is used as a NET, it identifies the network entity rather than a specific upper-layer service running on that entity. Therefore, the selector field must be set to 00. For example, a commonly configured NET such as 49.0001.0000.0000.0001.00 ends in 00; the preceding six bytes are the system ID and the final byte is the SEL.

This rule is intrinsic to IS-IS NET addressing and does not change according to whether the routing domain carries IPv4 routes, IPv6 routes, or both. IPv6 support adds IPv6 reachability information and related IS-IS extensions, but it does not redefine the NET format or the required selector value for the router's NET. The ISO IS-IS specification as incorporated for IP routing states that a NET is an NSAP address with the N-selector set to zero. Huawei IS-IS configurations consequently use NET values ending in .00. See [RFC 1195](#) and the IPv6 IS-IS extension specification in [RFC 5308](#).

QUESTION NO: 57

Regarding BGP MED, which of the following descriptions is correct?

- A. In BGP route selection, MED has a lower priority than AS_Path, Preferred-Value, Local_Pref, and Origin.
- B. The default value for the BGP routing MED is 0.
- C. By default, if a route has no MED attribute, MED is processed as 0. If the bestroute med-none-as-maximum command is configured, MED is processed as the maximum value, 4294967295.
- D. By default, BGP routing rules can compare the MED values of routes from different autonomous systems.

ANSWER: A B C

Explanation:

MED (Multi-Exit Discriminator) is an optional non-transitive BGP path attribute used to indicate a preferred entry point into an AS when multiple inter-AS links exist. In Huawei BGP best-path selection, a lower MED is preferred, but MED is evaluated only after higher-priority attributes such as Preferred-Value, Local_Pref, AS_Path length, and Origin. Therefore, its stated lower route-selection priority is correct. When a MED is not explicitly carried in a route, BGP normally treats the missing MED as 0, which is also the default MED behavior. Huawei VRP can alter that treatment with the `bestroute med-none-as-maximum` command: routes without a MED are then assigned the maximum unsigned 32-bit MED value, 4294967295, for best-route comparison. This permits routes that explicitly advertise a MED to be favored over routes for which MED is absent. The MED attribute format and the rule that lower values are preferred are defined in [RFC 4271](#); Huawei BGP configuration and command references are available through the [Huawei Enterprise Support Documentation portal](#).

QUESTION NO: 58

Which nationality meets the following two conditions: 1] BGP Router A can choose whether the Update message is based on the nationality. () If Router B receives a Update message containing a cross-hatch, run-by B does not have an identified solidity, but will also contain a multifaceted] Update retracts the advertisement to Router C. Router C may recognize and apply this attribute. (Multiple choice questions).

- A. Aggregator
- B. Local_Pref
- C. Multi_Exit_Disc
- D. Community

ANSWER: A D

Explanation:

The described behavior is that of an optional transitive BGP path attribute. Such an attribute can be included in BGP UPDATE messages, and a BGP speaker that does not recognize it must still propagate it to other peers. When forwarding an unrecognized optional transitive attribute, the speaker sets the Partial bit in the attribute flags, signaling that the attribute was not recognized by every AS through which the route passed. A later BGP speaker that supports the attribute can then recognize and process it.

Aggregator is an optional transitive attribute. It records the AS number and BGP identifier of the router that performed route aggregation, allowing this aggregation information to remain available as the route is advertised onward. **Community** is also an optional transitive attribute. Communities attach policy-significant tagging information to routes and are designed to be carried across AS boundaries, enabling routers that recognize a community value to apply routing policy based on that value.

These classifications and the required handling of unrecognized optional transitive attributes are specified in [RFC 4271, Border Gateway Protocol 4](#). The standardized Community attribute is further defined in [RFC 1997, BGP Communities Attribute](#).

QUESTION NO: 59

Configure queue-based traffic shaping, which defaults to tailing if the queue length exceeds the cache size

- A. True
- B. False

ANSWER: A

Explanation:

True is correct. In Huawei queue-based traffic shaping, packets are buffered when the configured shaping rate prevents them from being transmitted immediately. The queue therefore absorbs temporary bursts while traffic is regulated to the specified rate. If the queued traffic reaches and exceeds the available buffer (cache) capacity, the default packet-discard behavior is tail drop: subsequently arriving packets are discarded at the tail of the queue until sufficient buffer space becomes available. This protects the device from uncontrolled buffer consumption while allowing traffic shaping to continue enforcing the configured bandwidth limit. The behavior is important when dimensioning queue buffers and shaping parameters, because sustained offered traffic above the shaping rate can fill the queue and result in packet loss. Huawei QoS documentation describes traffic shaping as buffering excess traffic and transmitting it at the configured rate; its command references and QoS configuration material are available through the [Huawei Enterprise Documentation portal](#). The underlying differentiated-services architecture also recognizes queue management and packet dropping as mechanisms used when traffic exceeds available queue resources; see [RFC 2475](#).

QUESTION NO: 60

BGP EVPN does MAC address advertisement, ARP message advertisement, and IRB route advertisement for which BGP EVPN route is advertised?

- A. Type3
- B. Type2
- C. Type4
- D. Type1

ANSWER: B

Explanation:

Type2 is correct because the EVPN MAC/IP Advertisement route carries endpoint reachability information. In Huawei EVPN deployments, this route type is used to advertise a learned or configured MAC address and can also include the associated IP address. Advertising the IP-to-MAC binding enables distributed ARP/ND suppression: a receiving VTEP can answer an ARP request locally rather than flooding it across the EVPN fabric. The same route mechanism is also used for integrated routing and bridging reachability, allowing remote VTEPs to learn the host IP/MAC information needed to forward traffic correctly across Layer 2 and Layer 3 EVPN services. Thus, MAC advertisement, ARP-related host information advertisement, and IRB host-route advertisement are all functions associated with the MAC/IP Advertisement route. Huawei's EVPN documentation identifies the MAC/IP Advertisement route as EVPN route type 2 and describes its role in carrying MAC and IP address information for host reachability. See [Huawei EVPN Route Types](#) and [RFC 7432, MAC/IP Advertisement Route](#).

QUESTION NO: 61

As shown in the figure, there are two IPv6 networks that can access the IPv4 network, and the IPsec tunnel needs to be established between the two IPv6 networks to communicate, which of the following requirements?(). Single choice questions).



- A. ESP + tunnel mode
- B. AH + transmission mode

- C. AH + tunnel mode
- D. None of the above options are correct

ANSWER: A

Explanation:

ESP + tunnel mode is the appropriate IPsec protection method for establishing secure communication between IPv6 networks across an IPv4 transit network. Tunnel mode encapsulates the complete original IPv6 packet, including its IPv6 header, inside a newly formed IPsec packet. The resulting protected packet can use IPv4 as its outer header, allowing it to be routed through the intervening IPv4-only network while preserving the original IPv6 packet for delivery at the remote IPsec gateway.

ESP provides the protection normally required for a site-to-site VPN: payload confidentiality through encryption, integrity protection, data-origin authentication, and anti-replay protection when configured with suitable algorithms. At the receiving gateway, ESP processing removes the outer IPv4 and IPsec encapsulation and forwards the recovered native IPv6 packet into the destination IPv6 network. This is the normal design principle for a gateway-to-gateway VPN that protects traffic between separate address domains rather than only traffic originating from the gateways themselves.

The IPsec architecture defines tunnel mode as protection that encapsulates an entire inner IP packet, and the ESP specification describes its confidentiality and authentication capabilities. See [RFC 4301: Security Architecture for IP](#) and [RFC 4303: Encapsulating Security Payload](#).

QUESTION NO: 62

The following is the correct description of the aggregation of BGP? (Single Choice Questions).

- A. After you configure aggregate ipv4-address mask , only aggregate routes are published, not detail routes.
- B. For IPv6 routing, BGP supports manual coalescing and automatic aggregation
- C. By default, BGP enables automatic aggregation
- D. You can manually aggregate the routes in the BGP local routing table with a profile

ANSWER: D

Explanation:

You can manually aggregate the routes in the BGP local routing table with a profile. On Huawei VRP, manual BGP route aggregation is configured by creating an aggregate route for a specified IPv4 prefix and mask. BGP evaluates matching more-specific routes in its local BGP routing information and can originate the corresponding aggregate route. Aggregation is therefore an intentional policy action: it lets an operator advertise a summarized reachability statement rather than requiring every contributing prefix to be represented individually in external advertisements. Huawei's aggregation command also provides policy-related controls, allowing a route policy to influence aggregation behavior or attributes. This is the meaning of using a profile in the statement: the aggregation process can be associated with routing-policy-based control, rather than being an uncontrolled summarization mechanism. The aggregate is handled as BGP routing information and is advertised according to the normal BGP eligibility and export-policy process. This behavior aligns with the BGP specification's treatment of route aggregation, including the need to form and advertise an aggregate route deliberately and consistently. See Huawei's [Enterprise Support documentation portal](#) for the VRP BGP command reference, and the route-aggregation rules in [RFC 4271, section 9.2.2.2](#).

QUESTION NO: 63

What is the relationship between CIR, BC, and TC?

- A. CIR=TC/BC
- B. CIR=BE/TC

C. $TC = CIR/BC$

D. $TC = BC/CIR$

ANSWER: D

Explanation:

$TC = BC/CIR$ is correct because the committed burst size (BC) represents the quantity of committed traffic that may be sent during one traffic-control measurement interval (TC), while the committed information rate (CIR) is the committed transmission rate. Their fundamental relationship is $BC = CIR \times TC$. Rearranging this formula gives $TC = BC / CIR$. This relationship is used when configuring and validating token-bucket traffic policing or shaping behavior: for a selected rate and burst allowance, it establishes the time interval over which that burst corresponds to the configured committed rate. Units must be consistent. For example, if BC is expressed in bits and CIR in bit/s, TC is obtained in seconds; if the equipment uses bytes, kilobits, or milliseconds in its command syntax, the appropriate unit conversion must be made before calculating the value. Huawei traffic-policy documentation describes CAR parameters such as committed rate and committed burst size, and Huawei's QoS configuration guidance covers the rate-and-burst model used for traffic control. See [Huawei: Configuring CAR](#) and [Huawei: QoS Configuration](#).