

Google Cloud Certified - Professional Cloud Network Engineer

Google Professional-Cloud-Network-Engineer

Version Demo

Total Demo Questions: 30

Total Premium Questions: 306

Buy Premium PDF

<https://passqueen.com>

support@passqueen.com

Topic Break Down

Topic	No. of Questions
Topic 1, Designing, planning, and prototyping a Google Cloud network	77
Topic 2, Implementing a Google Cloud network	107
Topic 3, Managing, monitoring, and optimizing a Google Cloud network	46
Topic 4, Configuring network security, access, and policies	76
Total	306

QUESTION NO: 1

You have an application running on Compute Engine that uses BigQuery to generate some results that are stored in Cloud Storage. You want to ensure that none of the application instances have external IP addresses.

Which two methods can you use to accomplish this? (Choose two.)

- A. Enable Private Google Access on all the subnets.
- B. Enable Private Google Access on the VPC.
- C. Enable Private Services Access on the VPC.
- D. Create network peering between your VPC and BigQuery.
- E. Create a Cloud NAT, and route the application traffic via NAT gateway.

ANSWER: A E

Explanation:

Compute Engine VM instances without external IP addresses can access Google APIs and services, including BigQuery and Cloud Storage, by enabling Private Google Access on every subnet that contains those instances. Private Google Access is a subnet-level setting: it lets VMs that have only internal IP addresses reach Google APIs and services through Google's network, without assigning public IP addresses to the VMs. This is appropriate when the application's required outbound connectivity is limited to supported Google services.

Creating a Cloud NAT gateway is also a valid solution. Cloud NAT provides outbound connectivity for resources that use only internal IP addresses while preserving the absence of external IPs on individual VM instances. The NAT gateway performs source network address translation for egress traffic, so instances can initiate connections without being directly reachable from the internet. This approach supports required outbound access beyond private Google API connectivity as well. Both designs maintain the core requirement that application instances do not receive external IP addresses. See [Private Google Access documentation](#) and [Cloud NAT overview](#).

QUESTION NO: 2

Your company's current network architecture has two VPCs that are connected by a dual-NIC instance that acts as a bump-in-the-wire firewall between the two VPCs. Flows between pairs of subnets across the two VPCs are working correctly. Suddenly, you receive an alert that none of the flows between the two VPCs are working anymore. You need to troubleshoot the problem. What should you do? (Choose 2 answers)

- A. Verify that the dual-NIC instance has not been added to a backend service.
- B. Verify that a public IP address has not been assigned to any network interface of the dual-NIC instance.
- C. Use Cloud Logging to verify that there were no modifications to the VPC firewall rules or policies that were applied to the two network interfaces of the dual-NIC instance.
- D. Verify that a VPC Service Controls perimeter has not been enabled for the project that contains the two VPCs and the dual-NIC instance.
- E. Verify that the dual-NIC instance has the `--can-ip-forward` attribute enabled.

ANSWER: C E

Explanation:

Use Cloud Logging to verify that there were no modifications to the VPC firewall rules or policies that were applied to the two network interfaces of the dual-NIC instance. A sudden, broad failure of previously working flows can result from a configuration change that affects ingress or egress traffic at either network interface. Cloud Audit Logs provide an administrative-change record for Google Cloud resources, including changes to firewall rules and network firewall policies. Reviewing those logs for the relevant time window helps identify whether a rule, policy association, priority, target, or action was changed in a way that interrupts the transit path.

Verify that the dual-NIC instance has the `--can-ip-forward` attribute enabled. A VM functioning as a bump-in-the-wire firewall must receive and forward packets whose source or destination is not the VM itself. Google Cloud requires IP forwarding to be enabled on an instance that acts as a network appliance or router. This setting is an instance property, and its state should be confirmed while investigating loss of forwarding across both VPCs. See [Google Cloud documentation for IP forwarding](#) and [Cloud Audit Logs documentation](#).

QUESTION NO: 3

(You need to migrate multiple PostgreSQL databases from your on-premises data center to Google Cloud. You want to significantly improve the performance of your databases while minimizing changes to your data schema and application code. You expect to exceed 150 TB of data per geographical region. You want to follow Google-recommended practices and minimize your operational costs. What should you do?)

- A. Migrate your data to AlloyDB.
- B. Migrate your data to Spanner.
- C. Migrate your data to Firebase.
- D. Migrate your data to Bigtable.

ANSWER: B

Explanation:

Migrate your data to Spanner. The decisive requirement is more than 150 TB of data in each geographic region while retaining a relational database and minimizing operational work. Spanner is Google Cloud's fully managed, horizontally scalable relational database. It automatically shards and distributes data across compute resources, so it is designed to scale far beyond the storage capacity of a single traditional PostgreSQL deployment without requiring database sharding infrastructure to be managed by the customer. Its regional configuration provides strongly consistent relational transactions and high availability within a selected region, making it appropriate when data does not require multi-region replication.

Spanner's PostgreSQL interface provides PostgreSQL-compatible SQL syntax, common client-library support, and a familiar relational model. This helps reduce—though not necessarily eliminate—the schema and application migration work compared with moving to a non-relational data store. Google's migration guidance also positions the PostgreSQL interface as a path for applications that use PostgreSQL tooling and SQL to adopt Spanner's scalability. Because Spanner is fully managed, Google handles replication, maintenance, scaling, and availability operations, reducing the operational burden. For the specified data volume, its elastic distributed architecture is the appropriate fit. See the [Cloud Spanner overview](#) and [Spanner PostgreSQL interface documentation](#).

QUESTION NO: 4

You are in the process of deploying an internal HTTP(S) load balancer for your web server virtual machine (VM) Instances. What two prerequisite tasks must be completed before creating the load balancer?

Choose 2 answers

- A. Choose a region.
- B. Create firewall rules for health checks
- C. Reserve a static IP address for the load balancer
- D. Determine the subnet mask for a proxy-only subnet.
- E. Determine the subnet mask for Serverless VPC Access.

ANSWER: B C

Explanation:

Create firewall rules for health checks is required so that Google Cloud health-check probe traffic can reach the backend VM instances. Health checks determine whether a backend is eligible to receive traffic, so firewall policy must permit the applicable health-check source ranges and protocol/port used by the health check. For a regional internal Application Load Balancer, backend firewall policy must also account for traffic forwarded from the proxy-only subnet to the backends. This enables the load balancer to establish and maintain healthy backend capacity.

Reserve a static IP address for the load balancer provides the stable regional internal address used by clients to reach the internal HTTP(S) load balancer. Reserving the address before creating the forwarding rule lets the deployment explicitly associate the frontend with the intended IP address and preserves that address for DNS records, client configuration, and future load-balancer updates. The address must be compatible with the selected VPC network and region. Google Cloud's setup guidance covers the required health-check firewall access and frontend addressing as part of configuring an internal Application Load Balancer. See [Internal Application Load Balancer overview](#) and [Google Cloud health-check probe IP ranges](#).

QUESTION NO: 5

You are creating an instance group and need to create a new health check for HTTP(s) load balancing.

Which two methods can you use to accomplish this? (Choose two.)

- A. Create a new health check using the gcloud command line tool.
- B. Create a new health check using the VPC Network section in the GCP Console.
- C. Create a new health check, or select an existing one, when you complete the load balancer's backend configuration in the GCP Console.
- D. Create a new legacy health check using the gcloud command line tool.
- E. Create a new legacy health check using the Health checks section in the GCP Console.

ANSWER: A C

Explanation:

For HTTP(S) Load Balancing, you can create a health check directly with the Google Cloud CLI by using the appropriate `gcloud compute health-checks create` command, such as `gcloud compute health-checks create http` or `https`. This creates a non-legacy health check resource that can be attached to the backend service used by the load balancer. Health-check settings include the request protocol, port or named port, request path, check interval, timeout, and healthy or unhealthy thresholds.

You can also create a health check while configuring the load balancer's backend in the Google Cloud console. During backend configuration, the console provides a health-check selection control where an existing compatible health check can be selected or a new one can be created. This workflow associates the health check with the backend service so that the load balancer can determine which instance-group endpoints are eligible to receive traffic.

Google recommends using the current health-check resource types for modern proxy load balancers, including HTTP(S) Load Balancing. See the [Google Cloud health check concepts documentation](#) and the [gcloud compute health-checks create http reference](#).

QUESTION NO: 6

Your company has provisioned 2000 virtual machines (VMs) in the private subnet of your Virtual Private Cloud (VPC) in the us-east1 region. You need to configure each VM to have a minimum of 128 TCP connections to a public repository so that users can download software updates and packages over the internet. You need to implement a Cloud NAT gateway so that the VMs are able to perform outbound NAT to the internet. You must ensure that all VMs can simultaneously connect to the public repository and download software updates and packages. Which two methods can you use to accomplish this? (Choose two.)

- A. Configure the NAT gateway in manual allocation mode, allocate 2 NAT IP addresses, and update the minimum number of ports per VM to 256.

- B.** Create a second Cloud NAT gateway with the default minimum number of ports configured per VM to 64.
- C.** Use the default Cloud NAT gateway's NAT proxy to dynamically scale using a single NAT IP address.
- D.** Use the default Cloud NAT gateway to automatically scale to the required number of NAT IP addresses, and update the minimum number of ports per VM to 128.
- E.** Configure the NAT gateway in manual allocation mode, allocate 4 NAT IP addresses, and update the minimum number of ports per VM to 128.

ANSWER: D E

Explanation:

The required Cloud NAT capacity is 256,000 source ports: 2,000 VMs multiplied by the required minimum of 128 TCP connections per VM. A NAT IP address provides 64,512 usable TCP/UDP source ports for Cloud NAT port allocation. Therefore, configuring manual allocation with four NAT IP addresses and 128 minimum ports per VM supplies 258,048 available ports, which is sufficient for every VM to receive its configured minimum allocation simultaneously. This configuration also makes the public-IP capacity explicit and predictable.

Automatic NAT IP address allocation can also meet the requirement when the minimum number of ports per VM is set to 128. Cloud NAT calculates the needed IP-address capacity from the number of applicable VM network interfaces and their minimum port allocation, then automatically assigns sufficient NAT IP addresses from the configured auto-allocated address pool. With 2,000 VMs at 128 ports each, Cloud NAT allocates enough addresses to provide the required aggregate port capacity. The minimum port value guarantees the per-VM baseline, while automatic allocation manages the corresponding number of external addresses. See the [Cloud NAT ports and addresses documentation](#) and [Cloud NAT configuration documentation](#).

QUESTION NO: 7

Question:

Your organization's security team recently discovered that there is a high risk of malicious activities originating from some of your VMs connected to the internet. These malicious activities are currently undetected when TLS communication is used. You must ensure that encrypted traffic to the internet is inspected. What should you do?

- A.** Enable Cloud Armor TLS inspection policy, and associate the policy with the backend VMs.
- B.** Use Cloud NGFW Enterprise. Create a firewall rule for egress traffic with the `tls-inspect` flag and associate the firewall rules with the VMs.
- C.** Configure a TLS agent on every VM to intercept TLS traffic before it reaches the internet. Configure Sensitive Data Protection to analyze and allow/deny the content.
- D.** Use Cloud NGFW Essentials. Create a firewall rule for egress traffic and enable VPC Flow Logs with the TLS inspect option. Analyze the output logs content and block the outputs that have malicious activities.

ANSWER: B

Explanation:

Use Cloud NGFW Enterprise. Create a firewall rule for egress traffic with the `tls-inspect` flag and associate the firewall rules with the VMs. Cloud NGFW Enterprise supports TLS inspection for outbound connections, allowing the firewall to decrypt eligible TLS sessions, inspect their contents using configured security controls, and then re-encrypt the traffic before forwarding it to its internet destination. This provides visibility into malicious command-and-control traffic, malware delivery, and other threats that would otherwise be concealed by encryption.

The egress rule should be scoped to the affected workloads by using applicable firewall targeting, such as service accounts or network tags, and TLS inspection must be enabled on that rule. The implementation also requires the appropriate TLS inspection configuration and certificate authority setup so that inspected connections can be handled securely. This is a preventive, inline control: traffic can be evaluated and enforcement actions can be applied as it leaves the VM, rather than

relying only on post-event telemetry. Google documents TLS inspection as a Cloud NGFW Enterprise capability and provides configuration guidance in the [TLS inspection overview](#) and [TLS inspection configuration documentation](#).

QUESTION NO: 8

Your company is running out of network capacity to run a critical application in the on-premises data center. You want to migrate the application to GCP. You also want to ensure that the Security team does not lose their ability to monitor traffic to and from Compute Engine instances. Which two products should you incorporate into the solution? (Choose two.)

- A. VPC flow logs
- B. Firewall logs
- C. Cloud Audit logs
- D. Stackdriver Trace
- E. Compute Engine instance system logs

ANSWER: A B

Explanation:

VPC flow logs and **Firewall logs** provide complementary network-traffic visibility for Compute Engine workloads after migration. VPC Flow Logs record sampled network-flow metadata for VM interfaces in a subnet, including source and destination addresses, ports, protocol, bytes, packets, timestamps, and whether the connection was allowed or denied. This gives the Security team broad visibility into traffic entering and leaving instances, as well as traffic between workloads, and supports investigation, monitoring, and export to Cloud Logging, BigQuery, or SIEM tooling.

Firewall logs record connections evaluated by VPC firewall rules for rules on which logging is enabled. They capture the rule that applied and connection details, allowing security personnel to verify firewall-policy enforcement and investigate allowed or denied traffic at the control point. Enabling both capabilities preserves the two important perspectives typically needed during a data-center-to-cloud migration: observed network flows and firewall-rule decisions. Logging can be configured at the relevant subnet and firewall-rule scope, then retained or exported according to the organization's operational and security requirements. See Google Cloud's [VPC Flow Logs documentation](#) and [firewall rules logging documentation](#).

QUESTION NO: 9

Your Security team needs to capture copies of packets from selected Compute Engine instances and send them to a third-party network monitoring appliance. You want to use a managed Google Cloud packet-capture solution without changing application traffic paths.

Which two actions should you take? (Choose two.)

- A. Deploy the monitoring appliance instances as backends of an internal passthrough Network Load Balancer.
- B. Create a Packet Mirroring policy that uses the internal load balancer forwarding rule as the collector destination.
- C. Create a Cloud NAT gateway for the source instances.
- D. Create a VPC Network Peering connection between the source subnet and the monitoring subnet.
- E. Replace the default route with a custom route that uses the monitoring appliance as the next hop.

ANSWER: A B

Explanation:

Deploy the monitoring appliance instances as backends of an internal passthrough Network Load Balancer. Packet Mirroring sends mirrored packets to a collector forwarding rule, and the internal passthrough Network Load Balancer distributes those packets to the collector instances.

Create a Packet Mirroring policy that identifies the source instances, subnets, or tags to mirror and references the internal load balancer forwarding rule as the collector destination. Packet Mirroring copies traffic for observation; it does not inline the appliance or alter the original packet's forwarding path. See the [Packet Mirroring overview](#) and [Configure Packet Mirroring documentation](#).

QUESTION NO: 10

Your company has 10 separate Virtual Private Cloud (VPC) networks, with one VPC per project in a single region in Google Cloud. Your security team requires each VPC network to have private connectivity to the main on-premises location via a Partner Interconnect connection in the same region. To optimize cost and operations, the same connectivity must be shared with all projects. You must ensure that all traffic between different projects, on-premises locations, and the internet can be inspected using the same third-party appliances. What should you do?

- A.** Configure the third-party appliances with multiple interfaces and specific Partner Interconnect VLAN attachments per project. Create the relevant routes on the third-party appliances and VPC networks.
- B.** Configure the third-party appliances with multiple interfaces, with each interface connected to a separate VPC network. Create separate VPC networks for on-premises and internet connectivity. Create the relevant routes on the third-party appliances and VPC networks.
- C.** Consolidate all existing projects' subnetworks into a single VPC. Create separate VPC networks for on-premises and internet connectivity. Configure the third-party appliances with multiple interfaces, with each interface connected to a separate VPC network. Create the relevant routes on the third-party appliances and VPC networks.
- D.** Configure the third-party appliances with multiple interfaces. Create a hub VPC network for all projects, and create separate VPC networks for on-premises and internet connectivity. Create the relevant routes on the third-party appliances and VPC networks. Use VPC Network Peering to connect all projects' VPC networks to the hub VPC. Export custom routes from the hub VPC and import on all projects' VPC networks.

ANSWER: D

Explanation:

Configure the third-party appliances with multiple interfaces. Create a hub VPC network for all projects, and create separate VPC networks for on-premises and internet connectivity. Create the relevant routes on the third-party appliances and VPC networks. Use VPC Network Peering to connect all projects' VPC networks to the hub VPC. Export custom routes from the hub VPC and import on all projects' VPC networks. This design centralizes the inspection layer and the Partner Interconnect connectivity while preserving the existing project and VPC separation. The hub provides interfaces for the inspection appliances and connectivity domains, allowing traffic destined for on-premises networks, the internet, or other connected environments to be deliberately routed through the same appliance fleet. Each project VPC peers directly with the hub, so it can receive the hub's exported custom routes and send applicable traffic to the centralized inspection path. Route import and export must be enabled on the relevant peering connections, and route definitions must direct traffic to the appliance interfaces appropriately. This avoids duplicating Partner Interconnect VLAN attachments and inspection appliances across all ten projects, reducing both operational overhead and cost while retaining centralized policy enforcement. Google documents custom route exchange over VPC Network Peering and Partner Interconnect's private connectivity model in the [VPC Network Peering route exchange documentation](#) and the [Partner Interconnect overview](#).

QUESTION NO: 11

Your organization uses a Shared VPC host project. Application teams deploy VM instances in service projects, but the central network team must prevent application teams from assigning external IP addresses to newly created VM instances. Existing instances with external IP addresses must not be affected.

Which two actions should you take? (Choose two.)

- A.** Enforce the constraints/compute.vmExternallpAccess organization policy constraint.
- B.** Configure the policy to deny external IP access for applicable service projects, or allowlist only approved VM instances.
- C.** Create an ingress firewall rule that denies traffic from 0.0.0.0/0.

- D. Enable Private Google Access on all Shared VPC subnets.
- E. Remove the Compute Instance Admin role from all application-team users.

ANSWER: A B

Explanation:

Create an organization policy that enforces the `constraints/compute.vmExternalIpAccess` constraint. This constraint restricts which VM instances can have external IP addresses. It can be scoped at the organization, folder, or project level, allowing the organization to apply the restriction to the service projects where new application workloads are deployed.

Configure the policy with an allowlist that contains only approved VM instances, or configure it to deny external IP access for the applicable projects when no exceptions are required. Organization Policy is evaluated when resources are created or updated and provides preventive governance without requiring the network team to review every VM deployment manually. Existing resources are not automatically changed merely because a new policy is enforced. See the [Organization Policy constraints documentation](#) and the [restriction of external IP addresses documentation](#).

QUESTION NO: 12

In your company, two departments with separate GCP projects (code-dev and data-dev) in the same organization need to allow full cross-communication between all of their virtual machines in GCP. Each department has one VPC in its project and wants full control over their network. Neither department intends to recreate its existing computing resources. You want to implement a solution that minimizes cost.

Which two steps should you take? (Choose two.)

- A. Connect both projects using Cloud VPN.
- B. Connect the VPCs in project code-dev and data-dev using VPC Network Peering.
- C. Enable Shared VPC in one project (e. g., code-dev), and make the second project (e. g., data-dev) a service project.
- D. Enable firewall rules to allow all ingress traffic from all subnets of project code-dev to all instances in project data-dev, and vice versa.
- E. Create a route in the code-dev project to the destination prefixes in project data-dev and use nexthop as the default gateway, and vice versa.

ANSWER: B D

Explanation:

Connect the VPCs in project code-dev and data-dev using VPC Network Peering. provides private, internal IP connectivity between the two independently administered VPC networks. This preserves each department's ownership and control of its own project and VPC while allowing workloads to communicate without requiring existing VM resources to be recreated or moved. VPC Network Peering is designed for direct network-to-network connectivity in Google Cloud and does not require VPN gateways or public IP addressing. Google Cloud automatically exchanges subnet routes between peered networks, enabling instances to find the remote subnet ranges.

Enable firewall rules to allow all ingress traffic from all subnets of project code-dev to all instances in project data-dev, and vice versa. is required because VPC firewall rules are enforced independently in each VPC. Peering establishes the network path, but it does not override ingress firewall policy. Reciprocal ingress rules that use the peer VPC subnet CIDR ranges as sources allow the requested full bidirectional VM communication. Since the default implied egress rule permits outbound traffic unless it has been restricted, these ingress rules complete the connectivity configuration. See Google Cloud's [VPC Network Peering documentation](#) and [VPC firewall rules documentation](#).

QUESTION NO: 13

You have created an HTTP(S) load balanced service. You need to verify that your backend instances are responding properly.

How should you configure the health check?

- A. Set request-path to a specific URL used for health checking, and set proxy-header to PROXY_V1.
- B. Set request-path to a specific URL used for health checking, and set host to include a custom host header that identifies the health check.
- C. Set request-path to a specific URL used for health checking, and set response to a string that the backend service will always return in the response body.
- D. Set proxy-header to the default value, and set host to include a custom host header that identifies the health check.

ANSWER: C

Explanation:

Set request-path to a specific URL used for health checking, and set response to a string that the backend service will always return in the response body. This configuration gives the health check a purpose-built endpoint to query and an expected application-level response to validate. The request path should target a lightweight URL, such as `/healthz`, that can reliably report whether the application is able to handle requests. Configuring a response string causes the HTTP(S) health checker to look for that expected content in the backend response body, in addition to receiving a successful HTTP response. This makes the check more meaningful than simply confirming that a listener accepts connections or returns a generic success status: it verifies that the intended health-check handler and application response are available. The response body must consistently contain the configured string whenever the instance is healthy, so the endpoint should avoid dynamic or dependency-sensitive output unless those dependencies are deliberately part of the health criteria. Google Cloud documents the HTTP(S) health-check request path and expected response as configurable health-check properties. See [Health checks overview](#) and [gcloud compute health-checks create http reference](#).

QUESTION NO: 14

You are creating a new application and require access to Cloud SQL from VPC instances without public IP addresses.

Which two actions should you take? (Choose two.)

- A. Activate the Service Networking API in your project.
- B. Activate the Cloud Datastore API in your project.
- C. Create a private connection to a service producer.
- D. Create a custom static route to allow the traffic to reach the Cloud SQL API.
- E. Enable Private Google Access.

ANSWER: A C

Explanation:

To let VPC instances that do not have public IP addresses connect privately to a Cloud SQL instance, configure Cloud SQL private IP using private services access. This setup requires enabling the Service Networking API, which provides the Service Networking service used to establish and manage the private services access relationship between the consumer VPC network and Google-managed services such as Cloud SQL.

You must also create a private connection to a service producer. In practice, this means reserving an internal IP address range in the VPC and creating a private services access connection to `servicenetworking.googleapis.com`. Google uses addresses from that allocated range for the Cloud SQL instance's private IP. Workloads in the connected VPC can then send traffic directly to the instance's private address, without requiring a public IP address on either the Cloud SQL instance or the client VM.

After private services access is configured, create or update the Cloud SQL instance to use private IP and ensure applicable firewall and application settings permit the database connection. See the [Cloud SQL private IP configuration guide](#) and the [Private services access documentation](#).

QUESTION NO: 15

Question:

You are configuring the final elements of a migration effort where resources have been moved from on-premises to Google Cloud. While reviewing the deployed architecture, you noticed that DNS resolution is failing when queries are being sent to the on-premises environment. You log in to a Compute Engine instance, try to resolve an on-premises hostname, and the query fails. DNS queries are not arriving at the on-premises DNS server. You need to use managed services to reconfigure Cloud DNS to resolve the DNS error. What should you do?

- A.** Validate that the Compute Engine instances are using the Metadata Service IP address as their resolver. Configure an outbound forwarding zone for the on-premises domain pointing to the on-premises DNS server. Configure Cloud Router to advertise the Cloud DNS proxy range to the on-premises network.
- B.** Validate that there is network connectivity to the on-premises environment and that the Compute Engine instances can reach other on-premises resources. If errors persist, remove the VPC Network Peerings and recreate the peerings after validating the routes.
- C.** Review the existing Cloud DNS zones, and validate that there is a route in the VPC directing traffic destined to the IP address of the DNS servers. Recreate the existing DNS forwarding zones to forward all queries to the on-premises DNS servers.
- D.** Ensure that the operating systems of the Compute Engine instances are configured to send DNS queries to the on-premises DNS servers directly.

ANSWER: A

Explanation:

Validate that the Compute Engine instances are using the Metadata Service IP address as their resolver. Configure an outbound forwarding zone for the on-premises domain pointing to the on-premises DNS server. Configure Cloud Router to advertise the Cloud DNS proxy range to the on-premises network. This uses Cloud DNS's managed outbound forwarding capability to route queries for the specified on-premises DNS suffix to the designated on-premises DNS servers. Compute Engine VMs should send their DNS requests to the Google-provided metadata server resolver at `169.254.169.254`, which applies Cloud DNS policies and private-zone resolution behavior for the VPC network. For an outbound forwarding zone, Cloud DNS sends forwarded queries from source addresses in `35.199.192.0/19`. The on-premises environment must have a return route for that proxy range through the Cloud VPN tunnel or Cloud Interconnect attachment; otherwise, the DNS server cannot return responses to Cloud DNS. Advertising this range through Cloud Router over BGP provides the required dynamic on-premises route while retaining managed DNS resolution from the VM perspective. See [Cloud DNS forwarding zones documentation](#) and [Cloud DNS alternative name server and forwarding requirements](#).

QUESTION NO: 16

You need to define an address plan for a future new Google Kubernetes Engine (GKE) cluster in your Virtual Private Cloud (VPC). This will be a VPC-native cluster, and the default Pod IP range allocation will be used. You must pre-provision all the needed VPC subnets and their respective IP address ranges before cluster creation. The cluster will initially have a single node, but it will be scaled to a maximum of three nodes if necessary. You want to allocate the minimum number of Pod IP addresses. Which subnet mask should you use for the Pod IP address range?

- A.** /21
- B.** /22
- C.** /23
- D.** /25

ANSWER: B

Explanation:

/22 is the minimum Pod secondary IP range for this cluster. In a VPC-native GKE cluster using the default Pod IP address allocation, each node is configured for the default maximum of 110 Pods. GKE reserves a per-node Pod CIDR that is larger than the configured Pod maximum to support allocation behavior: for 110 Pods per node, GKE assigns a **/24** Pod CIDR per node, which contains 256 IP addresses. The planned maximum of three nodes therefore requires three /24 blocks, or 768 Pod IP addresses in total. A subnet range must be sized as a CIDR power-of-two block, so the smallest range that can contain at least 768 addresses is a /22, which provides 1,024 addresses. This range also leaves the required capacity available before the nodes are created or scaled, as required when pre-provisioning the secondary range. Google documents the relationship between maximum Pods per node, the per-node Pod CIDR, and the total Pod range required for VPC-native clusters in its [flexible Pod CIDR documentation](#) and [cluster sizing guidance](#).

QUESTION NO: 17

You created a new VPC network named Dev with a single subnet. You added a firewall rule for the network Dev to allow HTTP traffic only and enabled logging. When you try to log in to an instance in the subnet via Remote Desktop Protocol, the login fails. You look for the Firewall rules logs in Stackdriver Logging, but you do not see any entries for blocked traffic. You want to see the logs for blocked traffic.

What should you do?

- A. Check the VPC flow logs for the instance.
- B. Try connecting to the instance via SSH, and check the logs.
- C. Create a new firewall rule to allow traffic from port 22, and enable logs.
- D. Create a new firewall rule with priority 65500 to deny all traffic, and enable logs.

ANSWER: D

Explanation:

Create a new firewall rule with priority 65500 to deny all traffic, and enable logs. Firewall Rules Logging records connections only when they match a firewall rule for which logging has been explicitly enabled. The failed RDP connection does not match the HTTP allow rule, so it falls through to the VPC network's implied ingress deny rule. Although implied rules enforce the default deny behavior, logging cannot be enabled on them, which is why no blocked-traffic record appears.

An explicit ingress deny rule at priority 65500 is evaluated before the implied deny rule, whose priority is 65535. Therefore, traffic not otherwise permitted, including the attempted RDP connection, matches this explicit deny rule. Enabling logging on that rule generates Firewall Rules Logging entries for denied connections, including useful fields such as source and destination information, protocol, port, and the rule that made the decision. This provides visibility into the blocked RDP traffic while retaining the intended security posture of allowing only explicitly permitted ingress traffic. See Google Cloud's [Firewall Rules Logging documentation](#) and its [documentation on implied firewall rules](#).

QUESTION NO: 18

Your company has just launched a new critical revenue-generating web application. You deployed the application for scalability using managed instance groups, autoscaling, and a network load balancer as frontend. One day, you notice severe bursty traffic that caused autoscaling to reach the maximum number of instances, and users of your application cannot complete transactions. After an investigation, you think it is a DDOS attack. You want to quickly restore user access to your application and allow successful transactions while minimizing cost. Which two steps should you take? (Choose two.)

- A. Use Cloud Armor to blacklist the attacker's IP addresses.
- B. Increase the maximum autoscaling backend to accommodate the severe bursty traffic.
- C. Create a global HTTP(s) load balancer and move your application backend to this load balancer.

D. Shut down the entire application in GCP for a few hours. The attack will stop when the application is offline.

E. SSH into the backend compute engine instances, and view the auth logs and syslogs to further understand the nature of the attack.

ANSWER: A C

Explanation:

Use Cloud Armor to blacklist the attacker's IP addresses and create a global HTTP(S) load balancer and move your application backend to this load balancer. These measures work together to restore access while avoiding the cost of scaling application instances merely to absorb malicious requests. Cloud Armor is Google Cloud's web application firewall and DDoS-protection service for HTTP(S) applications. A security policy can include a deny rule matching known source IP addresses, causing those requests to be rejected at Google's edge before they consume load-balancer or backend capacity.

Cloud Armor policies are enforced with supported external HTTP(S) application load balancing; therefore, moving the web application from a network load balancer to a global HTTP(S) load balancer provides the required protected frontend. The global load balancer also gives the service a globally distributed anycast entry point and integrates the backend service with the Cloud Armor policy. This architecture filters identified attack traffic before it reaches the managed instance group, allowing legitimate transaction requests to use the existing backend capacity and autoscaling budget. IP-based blocking is particularly suitable when the investigation has identified attacker addresses and rapid mitigation is needed. Google recommends Cloud Armor with an external Application Load Balancer to protect internet-facing web workloads and apply edge security controls. See [Cloud Armor overview](#) and [External Application Load Balancer documentation](#).

QUESTION NO: 19

You have recently been put in charge of managing identity and access management for your organization. You have several projects and want to use scripting and automation wherever possible. You want to grant the editor role to a project member.

Which two methods can you use to accomplish this? (Choose two.)

A. `GetIamPolicy()` via REST API

B. `setIamPolicy()` via REST API

C. `gcloud pubsub add-iam-policy-binding Sprojectname --member user:Susername --role roles/editor`

D. `gcloud projects add-iam-policy-binding Sprojectname --member user:Susername --role roles/editor`

E. Enter an email address in the Add members field, and select the desired role from the drop-down menu in the GCP Console.

ANSWER: B D

Explanation:

`setIamPolicy()` via REST API is a supported programmatic method for changing a project's IAM policy. A script can first obtain the project's existing policy, add a binding that grants `roles/editor` to the required member, and submit the revised policy with `setIamPolicy`. When updating a policy, the policy version and etag should be preserved as appropriate so that the update does not unintentionally overwrite concurrent IAM changes. Google documents this method as the project resource API operation for setting an IAM policy.

`gcloud projects add-iam-policy-binding Sprojectname --member user:Susername --role roles/editor` is also intended for this task. The Google Cloud CLI command adds a member-to-role binding to the IAM policy of the specified project, making it suitable for shell scripts, CI/CD automation, and repeatable administrative workflows. The member value identifies the principal, such as a user email address prefixed by `user:`, and `roles/editor` is the predefined Editor role being granted. See the [projects.setIamPolicy REST API documentation](#) and the [gcloud projects add-iam-policy-binding reference](#).

QUESTION NO: 20

Question:

Your organization has a hub and spoke architecture with VPC Network Peering, and hybrid connectivity is centralized at the hub. The Cloud Router in the hub VPC is advertising subnet routes, but the on-premises router does not appear to be receiving any subnet routes from the VPC spokes. You need to resolve this issue. What should you do?

- A. Create custom learned routes at the Cloud Router in the hub to advertise the subnets of the VPC spokes.
- B. Create custom routes at the Cloud Router in the spokes to advertise the subnets of the VPC spokes.
- C. Create a BGP route policy at the Cloud Router, and ensure the subnets of the VPC spokes are being announced towards the on-premises environment.
- D. Configure custom route advertisements at the Cloud Router in the hub to advertise the subnets of the VPC spokes.

ANSWER: D

Explanation:

Configure custom route advertisements at the Cloud Router in the hub to advertise the subnets of the VPC spokes. Cloud Router automatically advertises subnet routes that belong to the VPC network where the router resides when it uses the default advertisement mode. However, a spoke subnet reached through VPC Network Peering belongs to a different VPC network, so it is not included automatically in the hub router's default BGP advertisements. This behavior prevents the on-premises BGP peer from learning the spoke prefixes even though the hub VPC itself has connectivity to those prefixes through peering.

Set the hub Cloud Router's advertisement mode to custom and explicitly add each required spoke subnet CIDR as a custom advertised IP range. The router then includes those prefixes in updates sent through its Cloud VPN or Cloud Interconnect BGP session. Because the hub already has VPC Network Peering routes to the spoke networks, it can forward traffic destined for the advertised spoke ranges after on-premises networks learn them. Ensure that the advertised ranges do not overlap with on-premises address space and that the corresponding VPC peering connectivity remains in place. See the Cloud Router [overview documentation](#) and the [VPC Network Peering documentation](#).

QUESTION NO: 21

You need to centralize the Identity and Access Management permissions and email distribution for the WebServices Team as efficiently as possible.

What should you do?

- A. Create a Google Group for the WebServices Team.
- B. Create a G Suite Domain for the WebServices Team.
- C. Create a new Cloud Identity Domain for the WebServices Team.
- D. Create a new Custom Role for all members of the WebServices Team.

ANSWER: A

Explanation:

Create a Google Group for the WebServices Team. Google Groups provides a single membership list that can be used both as an email distribution list and as a principal for Google Cloud IAM. Assigning IAM roles to the group, rather than individually to each team member, centralizes access administration: adding or removing a person from the group automatically grants or revokes the permissions that the group receives. This approach follows the Google Cloud security best practice of assigning roles to groups wherever practical, which reduces repetitive policy changes and makes access reviews easier. The same group email address can distribute messages to all members, satisfying the team communication requirement without maintaining a separate distribution mechanism. Group membership can be managed through Google Workspace or Cloud Identity, with appropriate controls for membership and ownership. For details, see [Using Google Groups with IAM](#) and [Learn about Google Groups](#).

QUESTION NO: 22

You work for a multinational enterprise that is moving to GCP.

These are the cloud requirements:

- An on-premises data center located in the United States in Oregon and New York with Dedicated Interconnects connected to Cloud regions us-west1 (primary HQ) and us-east4 (backup)
- Multiple regional offices in Europe and APAC
- Regional data processing is required in europe-west1 and australia-southeast1
- Centralized Network Administration Team

Your security and compliance team requires a virtual inline security appliance to perform L7 inspection for URL filtering. You want to deploy the appliance in us-west1.

What should you do?

- A.** • Create 2 VPCs in a Shared VPC Host Project. • Configure a 2-NIC instance in zone us-west1-a in the Host Project. • Attach NIC0 in VPC #1 us-west1 subnet of the Host Project. • Attach NIC1 in VPC #2 us-west1 subnet of the Host Project. • Deploy the instance. • Configure the necessary routes and firewall rules to pass traffic through the instance.
- B.** • Create 2 VPCs in a Shared VPC Host Project. • Configure a 2-NIC instance in zone us-west1-a in the Service Project. • Attach NIC0 in VPC #1 us-west1 subnet of the Host Project. • Attach NIC1 in VPC #2 us-west1 subnet of the Host Project. • Deploy the instance. • Configure the necessary routes and firewall rules to pass traffic through the instance.
- C.** • Create 1 VPC in a Shared VPC Host Project. • Configure a 2-NIC instance in zone us-west1-a in the Host Project. • Attach NIC0 in us-west1 subnet of the Host Project. • Attach NIC1 in us-west1 subnet of the Host Project. • Deploy the instance. • Configure the necessary routes and firewall rules to pass traffic through the instance.
- D.** • Create 1 VPC in a Shared VPC Service Project. • Configure a 2-NIC instance in zone us-west1-a in the Service Project. • Attach NIC0 in us-west1 subnet of the Service Project. • Attach NIC1 in us-west1 subnet of the Service Project. • Deploy the instance. • Configure the necessary routes and firewall rules to pass traffic through the instance.

ANSWER: B

Explanation:

Create 2 VPCs in a Shared VPC Host Project, then deploy the two-network-interface security appliance in a Service Project using subnets shared from the Host Project. Shared VPC centralizes ownership and administration of VPC networks, subnets, routes, firewall policies, and hybrid connectivity in the host project, while allowing separate service projects to deploy workloads. This separation is particularly appropriate for a centralized network team that needs to control enterprise networking while security or application teams operate their own compute resources.

A virtual inline appliance needs connectivity to distinct network paths so that traffic can be explicitly routed through it for inspection. A Compute Engine instance can have multiple network interfaces, with each interface connected to a different VPC network. For a VM deployed in a Shared VPC service project, its interfaces can use eligible subnets that belong to the Shared VPC host project. Placing both interfaces in us-west1 supports the requested appliance location, and custom routes plus firewall rules can steer the relevant traffic through the appliance for URL filtering.

See [Shared VPC documentation](#) and [multiple network interfaces documentation](#).

QUESTION NO: 23

You are configuring a Private Service Connect published service. Consumer VPC networks must access a single internal TCP service, but they must not receive general route access to the producer VPC. The service is deployed behind an internal passthrough Network Load Balancer.

Which two actions should you take? (Choose two.)

- A.** Create a Private Service Connect service attachment that references the internal passthrough Network Load Balancer forwarding rule.

- B. Configure a consumer acceptance policy on the service attachment.
- C. Create VPC Network Peering between every consumer VPC and the producer VPC.
- D. Configure Cloud NAT in the producer VPC for consumer access.
- E. Create an external Application Load Balancer with a public IP address for the service.

ANSWER: A B

Explanation:

Create a Private Service Connect service attachment that references the producer internal passthrough Network Load Balancer forwarding rule. The service attachment publishes the selected service so that consumers can create endpoints to it without requiring VPC Network Peering or broad route exchange.

Configure an appropriate consumer acceptance policy on the service attachment. A producer can automatically or manually accept connections from specified consumer projects or networks, which allows the producer to control which consumers can use the published service. Consumer endpoints then provide private access only to the exposed service. See the [Private Service Connect overview](#) and [publishing services by using Private Service Connect](#).

QUESTION NO: 24

You have ordered Dedicated Interconnect in the GCP Console and need to give the Letter of Authorization/Connecting Facility Assignment (LOA-CFA) to your cross-connect provider to complete the physical connection.

Which two actions can accomplish this? (Choose two.)

- A. Open a Cloud Support ticket under the Cloud Interconnect category.
- B. Download the LOA-CFA from the Hybrid Connectivity section of the GCP Console.
- C. Run `gcloud compute interconnects describe`.
- D. Check the email for the account of the NOC contact that you specified during the ordering process.
- E. Contact your cross-connect provider and inform them that Google automatically sent the LOA/CFA to them via email, and to complete the connection.

ANSWER: B D

Explanation:

"Download the LOA-CFA from the Hybrid Connectivity section of the GCP Console." is correct because, once Google has processed the Dedicated Interconnect order, the LOA-CFA is available in the console for the relevant interconnect. The document contains the information that the colocation facility or cross-connect provider needs to establish the physical cross-connect to Google's port. Downloading it from the console lets the customer deliver it to the provider through the provider's required process.

"Check the email for the account of the NOC contact that you specified during the ordering process." is also correct. Google sends LOA-CFA-related notification email to the technical/NOC contacts supplied with the Dedicated Interconnect order. Those contacts can access the authorization details needed to proceed with the physical connection. Keeping the designated NOC contact details current is therefore important during provisioning. Google's Dedicated Interconnect documentation describes obtaining the LOA-CFA and providing it to the colocation facility or cross-connect provider. See [Create Dedicated Interconnect connections](#) and [Dedicated Interconnect overview](#).

QUESTION NO: 25

You are configuring your Google Cloud environment to connect to your on-premises network. Your configuration must be able to reach Cloud Storage APIs and your Google Kubernetes Engine nodes across your private Cloud Interconnect

network. You have already configured a Cloud Router with your Interconnect VLAN attachments. You now need to set up the appropriate router advertisement configuration on the Cloud Router. What should you do?

- A. Configure the route advertisement to the default setting.
- B. On the on-premises router, configure a static route for the storage API virtual IP address which points to the Cloud Router's link-local IP address.
- C. Configure the route advertisement to the custom setting, and manually add prefix 199.36.153.8/30 to the list of advertisements. Leave all other options as their default settings.
- D. Configure the route advertisement to the custom setting, and manually add prefix 199.36.153.8/30 to the list of advertisements. Advertise all subnets visible to the Cloud Router.

ANSWER: D

Explanation:

Configure the route advertisement to the custom setting, manually add prefix 199.36.153.8/30 to the list of advertisements, and advertise all subnets visible to the Cloud Router. The 199.36.153.8/30 range is the Private Google Access VIP range for `private.googleapis.com`, which includes access to Google APIs such as Cloud Storage from on-premises hosts through Cloud Interconnect. Advertising this prefix through BGP gives the on-premises network a route for API traffic that is resolved to the private Google APIs VIP addresses. The corresponding private DNS configuration must resolve the applicable Google API names, such as Cloud Storage endpoints, to this VIP range. Because the router is being placed into custom advertisement mode to add the Google APIs prefix, it must also explicitly advertise the VPC subnet routes visible to the Cloud Router. This preserves on-premises reachability to the private IP addresses of GKE nodes in those VPC subnets. Cloud Router sends these selected prefixes to the BGP peer over the Interconnect VLAN attachment, allowing the on-premises network to learn both the Google APIs route and the node-subnet routes. See [Private Google Access for on-premises hosts](#) and [Cloud Router route advertisements](#).

QUESTION NO: 26

You have enabled HTTP(S) load balancing for your application, and your application developers have reported that HTTP(S) requests are not being distributed correctly to your Compute Engine Virtual Machine instances. You want to find data about how the request are being distributed.

Which two methods can accomplish this? (Choose two.)

- A. On the Load Balancer details page of the GCP Console, click on the Monitoring tab, select your backend service, and look at the graphs.
- B. In Stackdriver Error Reporting, look for any unacknowledged errors for the Cloud Load Balancers service.
- C. In Stackdriver Monitoring, select Resources > Metrics Explorer and search for `https/request_bytes_count` metric.
- D. In Stackdriver Monitoring, select Resources > Google Cloud Load Balancers and review the Key Metrics graphs in the dashboard.
- E. In Stackdriver Monitoring, create a new dashboard and track the `https/backend_request_count` metric for the load balancer.

ANSWER: A D

Explanation:

The Load Balancer details page's Monitoring tab is appropriate because it exposes monitoring charts for the selected backend service. These charts provide operational visibility into backend traffic and related serving behavior, allowing an engineer to inspect request activity for the service that routes traffic to the Compute Engine instances. Reviewing the backend-service-specific graphs is a direct way to investigate whether traffic is reaching the expected serving infrastructure.

The Google Cloud Load Balancers dashboard in Cloud Monitoring is also appropriate because its Key Metrics charts provide a centralized view of load-balancer telemetry, including request and backend-related measurements. These managed charts

are useful for correlating frontend demand with backend handling and for identifying traffic patterns while troubleshooting distribution concerns. Cloud Load Balancing automatically publishes metrics to Cloud Monitoring, so these views do not require application-side instrumentation. For current guidance on using the load balancer monitoring views and available metrics, see [Monitor load balancers](#) and [Google Cloud metrics](#).

QUESTION NO: 27

Your company has a single Virtual Private Cloud (VPC) network deployed in Google Cloud with on-premises connectivity already in place. You are deploying a new application using Google Kubernetes Engine (GKE), which must be accessible only from the same VPC network and on-premises locations. You must ensure that the GKE control plane is exposed to a predefined list of on-premises subnets through private connectivity only. What should you do?

- A.** Create a GKE private cluster with a private endpoint for the control plane. Configure VPC Networking Peering export/import routes and custom route advertisements on the Cloud Routers. Configure authorized networks to specify the desired on-premises subnets.
- B.** Create a GKE private cluster with a public endpoint for the control plane. Configure VPC Networking Peering export/import routes and custom route advertisements on the Cloud Routers.
- C.** Create a GKE private cluster with a private endpoint for the control plane. Configure authorized networks to specify the desired on-premises subnets.
- D.** Create a GKE public cluster. Configure authorized networks to specify the desired on-premises subnets.

ANSWER: A

Explanation:

Create a GKE private cluster with a private endpoint for the control plane. Configure VPC Networking Peering export/import routes and custom route advertisements on the Cloud Routers. Configure authorized networks to specify the desired on-premises subnets. A private endpoint ensures that control-plane access uses an internal IP address rather than the public internet, satisfying the requirement for private-only connectivity from the VPC and connected on-premises environment. The GKE control plane is reached through the cluster's VPC Network Peering connection, so the required routes must be exchanged appropriately between the customer VPC and the control-plane network. For hybrid connectivity, Cloud Router custom advertisements must also advertise the control-plane private IP range to the on-premises network, enabling return traffic and reliable reachability over Cloud VPN or Cloud Interconnect. Finally, authorized networks restrict which source CIDR ranges may connect to the control plane. Specifying only the approved on-premises subnets applies the required least-privilege access boundary while retaining private endpoint connectivity. See Google Cloud's guidance for [private GKE clusters](#) and its instructions for [authorized networks](#).

QUESTION NO: 28

You are planning to use Terraform to deploy the Google Cloud infrastructure for your company. The design must meet the following requirements

- Each Google Cloud project must represent an internal project that your team will work on
- After an internal project is finished, the infrastructure must be deleted
- Each internal project must have its own Google Cloud project owner to manage the Google Cloud resources.
- You have 10-100 projects deployed at a time

While you are writing the Terraform code, you need to ensure that the deployment is simple and the code is reusable with centralized management. What should you do?

- A.** Create a single project and additional VPCs for each internal project
- B.** Create a single shared VPC and attach each Google Cloud project as a service project
- C.** Create a single project and single VPC for each internal project

D. Create a Shared VPC and service project for each internal project

ANSWER: B

Explanation:

Create a Single Shared VPC and attach each Google Cloud project as a service project is correct because it separates each internal workload into its own service project while centralizing network administration. Terraform can use reusable modules to create each service project, assign that internal project's owner the required project-level IAM role, and deploy the project-specific resources. When an internal initiative ends, its service project and its associated Terraform-managed resources can be removed without requiring the shared network architecture to be recreated for every workload.

Shared VPC is specifically designed for this model: a centrally administered host project contains shared networking resources, while service projects consume those networks and retain separate identities, billing, IAM policy, quotas, and resource lifecycles. This is well suited to managing tens or hundreds of projects because the networking configuration is defined once in the host project and referenced consistently by reusable Terraform code. Project owners can manage appropriate resources in their individual service projects, while network administrators retain centralized control of subnets, routes, and firewall policy.

Google documents the host-project and service-project model in its [Shared VPC overview](#). For the Terraform implementation, use the Google provider's [google_project resource documentation](#) to provision the individual service projects consistently.

QUESTION NO: 29

Question:

You are configuring the firewall endpoints as part of the Cloud Next Generation Firewall (Cloud NGFW) intrusion prevention service in Google Cloud. You have configured a threat prevention security profile, and you now need to create an endpoint for traffic inspection. What should you do?

- A.** Attach the profile to the VPC network, create a firewall endpoint within the zone, and use a firewall policy rule to apply the L7 inspection.
- B.** Create a firewall endpoint within the zone, associate the endpoint to the VPC network, and use a firewall policy rule to apply the L7 inspection.
- C.** Create a firewall endpoint within the region, associate the endpoint to the VPC network, and use a firewall policy rule to apply the L7 inspection.
- D.** Create a Private Service Connect endpoint within the zone, associate the endpoint to the VPC network, and use a firewall policy rule to apply the L7 inspection.

ANSWER: B

Explanation:

Create a firewall endpoint within the zone, associate the endpoint to the VPC network, and use a firewall policy rule to apply the L7 inspection. Cloud NGFW firewall endpoints used for advanced traffic inspection are zonal resources. Creating the endpoint in the applicable zone establishes an inspection point for workloads whose traffic must be processed by Cloud NGFW in that zone. The endpoint must then be associated with the intended VPC network so that the network firewall policy can direct eligible traffic to it.

After the endpoint and network association exist, configure a network firewall policy rule for the required traffic direction, targets, protocols, and source or destination criteria. The rule uses the configured threat-prevention security profile (or applicable profile group) to enable Layer 7 intrusion-prevention inspection. This separates the inspection infrastructure—the zonal firewall endpoint—from enforcement logic in the firewall policy, allowing the policy to consistently apply the desired threat protections to matching traffic. For production resiliency and coverage, deploy endpoints in every zone containing workloads that require this inspection. See Google Cloud's [firewall endpoint configuration documentation](#) and [intrusion prevention configuration documentation](#).

QUESTION NO: 30

You have a VPC network with subnets in multiple regions. A Cloud Router in us-central1 learns on-premises routes through a Cloud Interconnect VLAN attachment. Workloads in europe-west1 must also use the learned routes to reach on-premises resources.

What should you do?

- A. Set the VPC network's dynamic routing mode to global.
- B. Create a VPC Network Peering connection between the us-central1 and europe-west1 subnets.
- C. Create a regional Cloud NAT gateway in europe-west1.
- D. Configure a default route in europe-west1 that uses the default internet gateway as its next hop.

ANSWER: A

Explanation:

Set the VPC network's dynamic routing mode to global. In global dynamic routing mode, routes learned by a Cloud Router can be used by eligible resources in all regions of the VPC network. This allows workloads in europe-west1 to use routes learned by the Cloud Router in us-central1.

In regional dynamic routing mode, dynamic routes are available only in the region that contains the Cloud Router that learned them. Changing the VPC to global dynamic routing provides the required cross-region availability while retaining the existing Cloud Router and VLAN attachment design. See the [Cloud Router dynamic routing mode documentation](#) and the [VPC routes overview](#).