

# **AWS Certified Solutions Architect - Associate (SAA-C03)**

**Amazon AWS SAA-C03**

**Version Demo**

**Total Demo Questions: 148**

**Total Premium Questions: 1483**

**Buy Premium PDF**

**<https://passqueen.com>**

**[support@passqueen.com](mailto:support@passqueen.com)**

## Topic Break Down

| Topic                                         | No. of Questions |
|-----------------------------------------------|------------------|
| Topic 1, Design Secure Architectures          | 483              |
| Topic 2, Design Resilient Architectures       | 323              |
| Topic 3, Design High-Performing Architectures | 387              |
| Topic 4, Design Cost-Optimized Architectures  | 285              |
| Topic 5, Mix Questions                        | 5                |
| Total                                         | 1483             |

#### QUESTION NO: 1

A company has a web server running on an Amazon EC2 instance in a public subnet with an Elastic IP address. The default security group is assigned to the EC2 instance. The default network ACL has been modified to block all traffic. A solutions architect needs to make the web server accessible from everywhere on port 443.

Which combination of steps will accomplish this task? (Choose two.)

- A. Create a security group with a rule to allow inbound TCP port 443 from source 0.0.0.0/0, and associate the security group with the EC2 instance.
- B. Create a security group with a rule to allow TCP port 443 to destination 0.0.0.0/0.
- C. Update the network ACL to allow TCP port 443 from source 0.0.0.0/0.
- D. Update the network ACL to allow inbound/outbound TCP port 443 from source 0.0.0.0/0 and to destination 0.0.0.0/0.
- E. Update the network ACL to allow inbound TCP port 443 from source 0.0.0.0/0 and outbound TCP port 1024-65535 to destination 0.0.0.0/0.

**ANSWER: A E**

#### Explanation:

Creating a security group that permits inbound TCP port 443 from 0.0.0.0/0 and associating it with the EC2 instance provides the required instance-level permission for HTTPS requests from any IPv4 address. Security groups are stateful: when an inbound HTTPS connection is allowed, response traffic for that connection is automatically permitted without requiring a separate outbound rule for the return traffic. The instance must actually be associated with this security group, because creating a security group alone does not change the rules applied to an existing instance.

Updating the network ACL with an inbound rule for TCP port 443 from 0.0.0.0/0 and an outbound ephemeral-port rule to 0.0.0.0/0 permits the request and its response at the subnet boundary. Unlike security groups, network ACLs are stateless, so explicit rules are required in both directions. HTTPS clients initiate connections to destination port 443, while the server's return packets are sent to the client ephemeral source port; allowing the ephemeral range enables those responses. Together, the security group and network ACL rules allow publicly initiated HTTPS connections to reach the Elastic IP and receive return traffic. See the [Amazon VPC security group documentation](#) and [network ACL documentation](#).

#### QUESTION NO: 2

A company is developing an ecommerce application that will consist of a load-balanced front end, a container-based application, and a relational database. A solutions architect needs to create a highly available solution that operates with as little manual intervention as possible.

Which solutions meet these requirements? (Select TWO.)

- A. Create an Amazon RDS DB instance in Multi-AZ mode.
- B. Create an Amazon RDS DB instance and one or more replicas in another Availability Zone.
- C. Create an Amazon EC2 in stance-based Docker cluster to handle the dynamic application load.
- D. Create an Amazon Elastic Container Service (Amazon ECS) cluster with a Fargate launch type to handle the dynamic application load.
- E. Create an Amazon Elastic Container Service (Amazon ECS) cluster with an Amazon EC2 launch type to handle the dynamic application load.

**ANSWER: A D**

#### Explanation:

Creating an Amazon RDS DB instance in Multi-AZ mode provides high availability for the relational database tier with minimal operational effort. RDS maintains a synchronous standby replica in a different Availability Zone and automatically

fails over to that standby if the primary DB instance has an infrastructure failure, Availability Zone disruption, or other qualifying outage. The application can continue using the database endpoint, avoiding the need to manually promote or reconfigure a replacement database instance. AWS manages the underlying replication and failover workflow. See the [Amazon RDS Multi-AZ documentation](#).

Creating an Amazon ECS cluster with a Fargate launch type supports the containerized application while minimizing infrastructure administration. AWS Fargate supplies and manages the compute capacity needed to run ECS tasks, so the company does not need to provision, patch, maintain, or replace Amazon EC2 container hosts. An ECS service can run tasks across multiple Availability Zones and integrate with load balancing; service deployment and scaling configuration can maintain the desired number of healthy tasks as demand changes or individual tasks fail. This combination provides an available, managed application tier suited to dynamic ecommerce workloads. See the [AWS Fargate for Amazon ECS documentation](#).

### QUESTION NO: 3

A company wants to create an API to authorize users by using JSON Web Tokens (JWTs). The company needs to support dynamic access to multiple AWS services by using path-based routing.

Which solution will meet these requirements?

- A. Deploy an Application Load Balancer behind an Amazon API Gateway REST API. Configure IAM authorization.
- B. Deploy an Application Load Balancer behind an Amazon API Gateway HTTP API. Use Amazon Cognito for authorization.
- C. Deploy a Network Load Balancer behind an Amazon API Gateway REST API. Use an AWS Lambda function as a custom authorizer.
- D. Deploy a Network Load Balancer behind an Amazon API Gateway HTTP API. Use Amazon Cognito for authorization.

### ANSWER: B

#### Explanation:

Deploy an Application Load Balancer behind an Amazon API Gateway HTTP API. Use Amazon Cognito for authorization. Amazon Cognito user pools authenticate users and issue JWTs, including ID and access tokens. An API Gateway HTTP API can validate JWTs from an OpenID Connect-compatible issuer such as a Cognito user pool by using a JWT authorizer. API Gateway validates the token signature and relevant claims, including the configured issuer and audience, before forwarding an authorized request to the backend integration. This provides managed, standards-based token authorization without requiring custom token-validation code.

An Application Load Balancer is the appropriate backend load balancer when requests must be distributed dynamically based on URL paths. Listener rules can match path patterns and forward requests to distinct target groups, allowing paths such as `/orders` and `/catalog` to reach separate services. API Gateway HTTP APIs support private integrations to an Application Load Balancer through a VPC link, combining JWT-protected API entry points with path-based service routing in the VPC. See the AWS documentation for [HTTP API JWT authorizers](#) and [Application Load Balancer listener rules and path conditions](#).

### QUESTION NO: 4

A company recently launched a variety of new workloads on Amazon EC2 instances in its AWS account. The company needs to create a strategy to access and administer the instances remotely and securely. The company needs to implement a repeatable process that works with native AWS services and follows the AWS Well-Architected Framework.

Which solution will meet these requirements with the LEAST operational overhead?

- A. Use the EC2 serial console to directly access the terminal interface of each instance for administration.
- B. Attach the appropriate IAM role to each existing instance and new instance. Use AWS Systems Manager Session Manager to establish a remote administrative session.

C. Create an administrative SSH key pair. Load the public key into each EC2 instance. Deploy a bastion host in a public subnet to provide a tunnel for administration of each instance.

D. Establish an AWS Site-to-Site VPN connection. Instruct administrators to use their local on-premises machines to connect directly to the instances by using SSH keys across the VPN tunnel.

**ANSWER: B**

**Explanation:**

Attach the appropriate IAM role to each existing instance and new instance. Use AWS Systems Manager Session Manager to establish a remote administrative session. This approach uses a native AWS managed service to provide secure, auditable access to EC2 instances without requiring inbound SSH access, public IP addresses, or administrator-managed bastion infrastructure. Instances need the Systems Manager Agent and an instance profile that grants the required Systems Manager managed-instance permissions, commonly through the `AmazonSSMManagedInstanceCore` policy. Administrators authenticate through IAM, and Session Manager can record session activity to Amazon S3 or Amazon CloudWatch Logs, supporting operational visibility and security governance.

The process is repeatable because the instance profile and Systems Manager Agent configuration can be incorporated into launch templates, Auto Scaling groups, images, or infrastructure-as-code deployments. Existing instances can receive the same IAM role configuration. Session Manager aligns with Well-Architected security practices by reducing the attack surface and replacing long-lived SSH key distribution with centralized IAM-based authorization. For private instances, Systems Manager connectivity can use outbound internet access or interface VPC endpoints, keeping management traffic within AWS networking where required. AWS documents Session Manager as a way to manage instances without opening inbound ports or maintaining bastion hosts. See [AWS Systems Manager Session Manager](#) and [Configure instance permissions for Systems Manager](#).

**QUESTION NO: 5**

A software company needs to upgrade a critical web application. The application is hosted in a public subnet. The EC2 instance runs a MySQL database. The application's DNS records are published in an Amazon Route 53 zone.

A solutions architect must reconfigure the application to be scalable and highly available. The solutions architect must also reduce MySQL read latency.

Which combination of solutions will meet these requirements? (Select TWO.)

A. Launch a second EC2 instance in a second AWS Region. Use a Route 53 failover routing policy to redirect the traffic to the second EC2 instance.

B. Create and configure an Auto Scaling group to launch private EC2 instances in multiple Availability Zones. Add the instances to a target group behind a new Application Load Balancer.

C. Migrate the database to an Amazon Aurora MySQL cluster. Create the primary DB instance and reader DB instance in separate Availability Zones.

D. Create and configure an Auto Scaling group to launch private EC2 instances in multiple AWS Regions. Add the instances to a target group behind a new Application Load Balancer.

E. Migrate the database to an Amazon Aurora MySQL cluster with cross-Region read replicas.

**ANSWER: B C**

**Explanation:**

**Create and configure an Auto Scaling group to launch private EC2 instances in multiple Availability Zones. Add the instances to a target group behind a new Application Load Balancer.** provides a resilient, scalable web tier. An Application Load Balancer distributes requests only to healthy registered targets, while an Auto Scaling group can maintain and adjust application capacity. Placing the instances in private subnets across multiple Availability Zones removes direct public exposure of the application servers and allows the load balancer to continue serving traffic if an Availability Zone or individual instance becomes unavailable. Route 53 can direct the application DNS record to the load balancer.

**Migrate the database to an Amazon Aurora MySQL cluster. Create the primary DB instance and reader DB instance in separate Availability Zones.** provides a highly available MySQL-compatible database architecture and a dedicated read scaling path. Aurora supports read replicas and exposes a reader endpoint that load balances connections across available Aurora Replicas. The application can send read-only workload to that endpoint, reducing read pressure and latency at the writer while retaining the writer endpoint for transactions that require writes. Deploying the writer and reader in separate Availability Zones also improves availability for the database tier. See the [Amazon EC2 Auto Scaling benefits documentation](#) and the [Amazon Aurora documentation](#).

#### QUESTION NO: 6

A company hosts an application in an Amazon EC2 Auto Scaling group. The company has observed that during periods of high demand, new instances take too long to join the Auto Scaling group and serve the increased demand. The company determines that the root cause of the issue is the long boot time of the instances in the Auto Scaling group. The company needs to reduce the time required to launch new instances to respond to demand. Which solution will meet this requirement?

- A. Increase the maximum capacity of the Auto Scaling group by 50%.
- B. Create a warm pool for the Auto Scaling group. Use the default specification for the warm pool size.
- C. Increase the health check grace period for the Auto Scaling group by 50%.
- D. Create a scheduled scaling action. Set the desired capacity equal to the maximum capacity of the Auto Scaling group.

**ANSWER: B**

#### Explanation:

Create a warm pool for the Auto Scaling group. Use the default specification for the warm pool size. A warm pool maintains pre-initialized EC2 instances that are ready to be placed into service when the Auto Scaling group scales out. Instead of waiting for a newly launched instance to complete the full provisioning and bootstrapping process during a demand spike, Auto Scaling can move an already prepared warm-pool instance into the active group. This substantially reduces the time before the instance can begin handling application traffic.

Warm-pool instances can be kept in a stopped, running, or hibernated state, depending on the selected pool state and application requirements. The default configuration is suitable for this requirement because it establishes an automatically managed reserve of prepared capacity without requiring the company to calculate a custom pool size first. Auto Scaling replenishes the warm pool after instances are promoted to the active group, preserving faster response to later scaling events. This feature is specifically intended for workloads with lengthy instance initialization, such as applications that install software, retrieve configuration, or perform other startup tasks before becoming usable. AWS documents warm pools and their lifecycle behavior at [Amazon EC2 Auto Scaling warm pools](#) and explains their use for reducing scale-out latency in the [Amazon EC2 Auto Scaling benefits documentation](#).

#### QUESTION NO: 7

A company has an Amazon S3 bucket that contains critical data. The company must protect the data from accidental deletion.

Which combination of steps should a solutions architect take to meet these requirements? (Choose two.)

- A. Enable versioning on the S3 bucket.
- B. Enable MFA Delete on the S3 bucket.
- C. Create a bucket policy on the S3 bucket.
- D. Enable default encryption on the S3 bucket.
- E. Create a lifecycle policy for the objects in the S3 bucket.

**ANSWER: A B**

### Explanation:

Enable versioning on the S3 bucket and enable MFA Delete on the S3 bucket. S3 Versioning preserves every version of an object stored in the bucket. When a user deletes an object in a versioning-enabled bucket, Amazon S3 adds a delete marker rather than permanently removing the prior object version. The previous version can therefore be recovered by removing the delete marker or restoring the required version. This provides the essential recovery mechanism for accidental object deletion.

MFA Delete adds an additional safeguard for the most consequential versioning operations. When MFA Delete is enabled, permanently deleting an object version or changing the bucket's versioning state requires authentication using the root user's MFA device. Requiring an MFA code helps prevent an accidental command, compromised credentials, or an unauthorized administrative action from permanently removing recoverable data. MFA Delete is configured through the S3 API or CLI and must be enabled by the bucket owner's root account. Together, versioning provides retention of prior object versions, while MFA Delete protects those versions from irreversible deletion. For implementation details, see the [Amazon S3 Versioning documentation](#) and [Amazon S3 MFA Delete documentation](#).

### QUESTION NO: 8

A company wants to migrate its on-premises data center to AWS. According to the company's compliance requirements, the company can use only the ap-northeast-3 Region. Company administrators are not permitted to connect VPCs to the internet.

Which solutions will meet these requirements? (Choose two.)

- A.** Use AWS Control Tower to implement data residency guardrails to deny internet access and deny access to all AWS Regions except ap-northeast-3.
- B.** Use rules in AWS WAF to prevent internet access. Deny access to all AWS Regions except ap-northeast-3 in the AWS account settings.
- C.** Use AWS Organizations to configure service control policies (SCPs) that prevent VPCs from gaining internet access. Deny access to all AWS Regions except ap-northeast-3.
- D.** Create an outbound rule for the network ACL in each VPC to deny all traffic from 0.0.0.0/0. Create an IAM policy for each user to prevent the use of any AWS Region other than ap-northeast-3.
- E.** Use AWS Config to activate managed rules to detect and alert for internet gateways and to detect and alert for new resources deployed outside of ap-northeast-3.

### ANSWER: A C

### Explanation:

Using AWS Control Tower to apply data-residency and internet-access guardrails provides centralized governance across accounts in a landing zone. Control Tower preventive controls can enforce restrictions through policies, while Region governance can restrict AWS operations to the approved Region. This provides an account-governance approach rather than relying on individual administrators to consistently configure each VPC or resource correctly. AWS Control Tower documents its controls and Region-deny capabilities at [AWS Control Tower controls](#).

Using AWS Organizations service control policies (SCPs) is also an appropriate centralized preventative solution. An SCP can explicitly deny EC2 API actions that would create or attach internet gateways, preventing administrators in member accounts from establishing VPC internet connectivity even when their IAM permissions would otherwise allow it. An SCP can also use the `aws:RequestedRegion` global condition key to deny requests outside `ap-northeast-3`. SCPs set the maximum permissions available to principals in organization member accounts, making these restrictions consistently enforceable across the environment. AWS provides examples of Region-restriction SCPs in [Service control policy examples](#).

### QUESTION NO: 9

A solutions architect is designing a two-tier web application. The application consists of a public-facing web tier hosted on Amazon EC2 in public subnets. The database tier consists of Microsoft SQL Server running on Amazon EC2 in a private subnet. Security is a high priority for the company.

How should security groups be configured in this situation? (Select TWO )

- A. Configure the security group for the web tier to allow inbound traffic on port 443 from 0.0.0.0/0.
- B. Configure the security group for the web tier to allow outbound traffic on port 443 from 0.0.0.0/0.
- C. Configure the security group for the database tier to allow inbound traffic on port 1433 from the security group for the web tier.
- D. Configure the security group for the database tier to allow outbound traffic on ports 443 and 1433 to the security group for the web tier.
- E. Configure the security group for the database tier to allow inbound traffic on ports 443 and 1433 from the security group for the web tier.

**ANSWER: A C**

**Explanation:**

Configure the security group for the web tier to allow inbound traffic on port 443 from 0.0.0.0/0. This allows users on the internet to reach the public-facing application over HTTPS, while limiting externally initiated access to the encrypted web application protocol. Because the web-tier instances are the intended internet entry point, the rule belongs on the security group associated with those instances. Security groups are stateful, so return traffic for permitted inbound HTTPS connections is automatically allowed without requiring a corresponding outbound rule for ephemeral response traffic.

Configure the security group for the database tier to allow inbound traffic on port 1433 from the security group for the web tier. Port 1433 is the standard SQL Server listener port, and using the web tier's security group as the source restricts database connectivity to instances that are members of that application tier. This implements tier-to-tier access control without relying on individual private IP addresses, which can change when EC2 instances are replaced or scaled. The database remains private and accepts only the specific application traffic required for the web tier to communicate with SQL Server. AWS security-group rules can reference another security group to authorize traffic from associated resources. See [AWS security groups](#) and [Security group rules](#).

**QUESTION NO: 10**

A company is running a multi-tier web application on premises. The web application is containerized and runs on a number of Linux hosts connected to a PostgreSQL database that contains user records. The operational overhead of maintaining the infrastructure and capacity planning is limiting the company's growth. A solutions architect must improve the application's infrastructure.

Which combination of actions should the solutions architect take to accomplish this? (Choose two.)

- A. Migrate the PostgreSQL database to Amazon Aurora
- B. Migrate the web application to be hosted on Amazon EC2 instances.
- C. Set up an Amazon CloudFront distribution for the web application content.
- D. Set up Amazon ElastiCache between the web application and the PostgreSQL database.
- E. Migrate the web application to be hosted on AWS Fargate with Amazon Elastic Container Service (Amazon ECS).

**ANSWER: A E**

**Explanation:**

**Migrate the PostgreSQL database to Amazon Aurora** is appropriate because Amazon Aurora provides a fully managed relational database service and offers PostgreSQL compatibility. This enables the company to retain PostgreSQL-oriented application and database tooling while offloading much of the operational work associated with database infrastructure, including provisioning, patching, backups, failure detection, and storage management. Aurora also supports automatic storage growth and can provide high availability through replication across Availability Zones, reducing both administration and capacity-planning effort. AWS documents Aurora PostgreSQL compatibility and managed database capabilities at [Amazon Aurora overview](#).

**Migrate the web application to be hosted on AWS Fargate with Amazon Elastic Container Service (Amazon ECS).** directly addresses the operational burden of operating Linux container hosts. Fargate is a serverless compute engine for containers: the company defines container images, task requirements, and scaling behavior, while AWS provisions and manages the underlying compute infrastructure. ECS can scale the number of running tasks in response to demand, avoiding the need to purchase, maintain, patch, and capacity-plan hosts. This preserves the existing containerized application architecture while making the platform more elastic and operationally efficient. See [AWS Fargate for Amazon ECS](#) for details.

#### QUESTION NO: 11

A company must migrate 20 TB of data from a data center to the AWS Cloud within 30 days. The company's network bandwidth is limited to 15 Mbps and cannot exceed 70% utilization. What should a solutions architect do to meet these requirements?

- A. Use AWS Snowball.
- B. Use AWS DataSync.
- C. Use a secure VPN connection.
- D. Use Amazon S3 Transfer Acceleration.

**ANSWER: A**

#### Explanation:

Use AWS Snowball. The usable network bandwidth is limited to 70% of 15 Mbps, or 10.5 Mbps. At that rate, transferring 20 TB over the network would take substantially longer than 30 days, even before accounting for protocol overhead, retries, or other network traffic. A physical data-transfer device is therefore the appropriate approach because it avoids making the migration dependent on the constrained WAN connection.

AWS Snowball devices are designed for secure, offline transfer of large datasets into AWS. The company copies the data from its data center to the device over its local network, ships the device to AWS, and AWS imports the data into the selected Amazon S3 bucket. This workflow supports terabyte-scale migrations while meeting a short migration window that the available internet bandwidth cannot satisfy. Snowball also uses encryption and tamper-resistant hardware to protect data during transit, and AWS manages device logistics and the import process. For service details and supported data-transfer workflows, see the [AWS Snowball product page](#) and the [AWS Snowball Edge Developer Guide](#).

#### QUESTION NO: 12

A company runs an application as a task in an Amazon Elastic Container Service (Amazon ECS) cluster. The application must have read and write access to a specific group of Amazon S3 buckets. The S3 buckets are in the same AWS Region and AWS account as the ECS cluster. The company needs to grant the application access to the S3 buckets according to the principle of least privilege.

Which combination of solutions will meet these requirements? (Select TWO.)

- A. Add a tag to each bucket. Create an IAM policy that includes a StringEquals condition that matches the tags and values of the buckets.
- B. Create an IAM policy that lists the full Amazon Resource Name (ARN) for each S3 bucket.
- C. Attach the IAM policy to the instance role of the ECS task.
- D. Create an IAM policy that includes a wildcard Amazon Resource Name (ARN) that matches all combinations of the S3 bucket names.
- E. Attach the IAM policy to the task role of the ECS task.

**ANSWER: B E**

### Explanation:

Creating an IAM policy that lists the full Amazon Resource Name (ARN) for each S3 bucket lets the company explicitly restrict access to only the required buckets. For S3 read and write access, the policy should scope bucket-level permissions, such as listing a bucket, to each bucket ARN and object-level permissions, such as `s3:GetObject` and `s3:PutObject`, to the corresponding object ARNs (for example, `arn:aws:s3:::example-bucket/*`). Explicit resource ARNs provide the precise resource boundary required for least privilege.

Attaching the IAM policy to the task role of the ECS task makes those permissions available through the task's AWS credentials. ECS task roles are intended to provide application-specific AWS permissions to containers, separating application permissions from the permissions used by the ECS container instance or ECS infrastructure. This enables the particular application task to access its approved S3 resources without extending those permissions to unrelated workloads. The task definition must specify the task role, and the role's trust policy must allow ECS tasks to assume it. AWS documents task roles as the mechanism for granting containers access to AWS services and recommends granting only the permissions the application needs. See [Amazon ECS task IAM roles](#) and [Amazon S3 policy resources and actions](#).

### QUESTION NO: 13

A company has an application that generates a large number of files. File sizes vary between 1 KB and 10 MB. The company wants to build a store-and-forward solution that keeps the files in a temporary Amazon S3 location. The solution must then preprocess the files and archive them to S3 Glacier Flexible Retrieval.

Which combination of steps will meet these requirements MOST cost-effectively? (Select THREE.)

- A. Configure the temporary S3 location to use the S3 Standard storage class.
- B. Use S3 Event Notifications to invoke an AWS Lambda function when a new file is moved into the temporary S3 location.
- C. Use event source mappings to invoke an AWS Lambda function when a new file is moved into the temporary S3 location.
- D. For files that are 128 KB or smaller, use an AWS Lambda function to add the files to a ZIP archive. Move the ZIP archive to S3 Glacier Flexible Retrieval when the ZIP file reaches 1 MB.
- E. For files that are 128 KB or smaller, use an AWS Lambda function to change the files' storage class to S3 One Zone-Infrequent Access (S3 One Zone-IA).
- F. Configure the temporary S3 location to use the S3 Standard-Infrequent Access (S3 Standard-IA) storage class.

### ANSWER: A B D

### Explanation:

Configuring the temporary S3 location to use the S3 Standard storage class provides a durable, highly available staging area without a minimum storage-duration charge. This is well suited to a store-and-forward workflow in which objects are expected to be processed and moved soon after upload. Amazon S3 Event Notifications can directly invoke an AWS Lambda function when a new object is created in the temporary bucket, creating an event-driven preprocessing workflow without polling or separately managing event consumers.

For files that are 128 KB or smaller, using AWS Lambda to add files to a ZIP archive and moving that archive to S3 Glacier Flexible Retrieval when it reaches 1 MB improves archival economics by reducing the number of individually archived objects and associated per-object overhead. S3 Glacier Flexible Retrieval charges include archival-object metadata, with part of that metadata stored and charged at S3 Standard rates. Aggregating numerous small files therefore reduces metadata and request overhead while preserving the required archival destination. The ZIP archive can be restored and processed as a single archive when access is needed. See the [Amazon S3 Event Notifications documentation](#) and [Amazon S3 storage class documentation](#).

### QUESTION NO: 14

A company needs to design a hybrid network architecture. The company's workloads are currently stored in the AWS Cloud and in on-premises data centers. The workloads require single-digit latencies to communicate. The company uses an AWS Transit Gateway transit gateway to connect multiple VPCs.

Which combination of steps will meet these requirements MOST cost-effectively? (Select TWO.)

- A. Establish an AWS Site-to-Site VPN connection to each VPC.
- B. Associate an AWS Direct Connect gateway with the transit gateway that is attached to the VPCs.
- C. Establish an AWS Site-to-Site VPN connection to an AWS Direct Connect gateway.
- D. Establish an AWS Direct Connect connection. Create a transit virtual interface (VIF) to a Direct Connect gateway.
- E. Associate AWS Site-to-Site VPN connections with the transit gateway that is attached to the VPCs

**ANSWER: B D**

**Explanation:**

**Establish an AWS Direct Connect connection. Create a transit virtual interface (VIF) to a Direct Connect gateway.** provides the private, dedicated connectivity path between the on-premises environment and AWS. A transit VIF is specifically designed to connect a Direct Connect connection to an AWS Direct Connect gateway. This architecture provides consistent network performance and avoids sending hybrid traffic over the public internet, supporting the low-latency requirement when the Direct Connect location and network path can provide the needed latency.

**Associate an AWS Direct Connect gateway with the transit gateway that is attached to the VPCs.** extends that single Direct Connect path to the VPCs connected through the Transit Gateway. The Direct Connect gateway acts as the association point between the transit VIF and the Transit Gateway, enabling the on-premises network to reach multiple attached VPCs without requiring separate dedicated connections for each VPC. This centralizes hybrid connectivity and reduces operational overhead while making efficient use of the existing Transit Gateway design. AWS documents that a transit VIF connects to a Direct Connect gateway, and that Direct Connect gateways can be associated with transit gateways for this connectivity model. See [AWS Direct Connect gateways and transit gateways](#) and [AWS virtual interfaces documentation](#).

## QUESTION NO: 15

A company needs a solution to give customers the ability to upload encrypted files to a directory in an Amazon S3 bucket by using SFTP. After customers upload files, the solution must automatically decrypt the files and move them to a second directory within the same S3 bucket for downstream processing.

The solution must not require authentication services. The solution must fully automate all post-upload operations and require minimal ongoing operational overhead.

Which solution will meet these requirements? (Select THREE.)

- A. Use AWS Transfer Family with the SFTP protocol. Configure the S3 bucket as the home directory for uploaded files.
- B. Use an S3 event notification to invoke an AWS Lambda function that moves uploaded files between folders.
- C. Use an AWS Transfer Family workflow and a DECRYPT action to decrypt uploaded files.
- D. Tag incoming S3 objects. Periodically query objects by using an external script that runs in a container.
- E. Use an AWS Transfer Family workflow and a COPY action to move files to a new directory within the S3 bucket after decryption.
- F. Use an AWS Batch job to poll the S3 bucket and run a decryption script on new files.

**ANSWER: A C E**

**Explanation:**

AWS Transfer Family provides a fully managed SFTP endpoint that can use Amazon S3 as its storage backend. Configuring the S3 bucket as the home directory gives customers a managed way to upload their encrypted files directly into the designated S3 location, without the company operating SFTP servers. Transfer Family supports service-managed users, so the implementation does not need a separate external authentication service.

A managed workflow associated with the Transfer Family server runs automatically after a successful upload. The workflow's `DECRYPT` step performs supported PGP decryption as a managed post-upload operation, avoiding the need to maintain decryption code, polling infrastructure, or a separate processing fleet. A subsequent `COPY` step can write the processed output to the required destination prefix in the same S3 bucket, making it available to downstream processing. These workflow steps can be sequenced, which provides an automated managed path from upload through decryption and placement in the processing directory. AWS documents Transfer Family SFTP and S3 integration at [AWS Transfer Family server with Amazon S3 storage](#) and describes the available managed workflow processing steps at [AWS Transfer Family workflow steps](#).

#### QUESTION NO: 16

A company stores sensitive customer data in an Amazon DynamoDB table. The company frequently updates the data. The company wants to use the data to personalize offers for customers.

The company's analytics team has its own AWS account. The analytics team runs an application on Amazon EC2 instances that needs to process data from the DynamoDB tables. The company needs to follow security best practices to create a process to regularly share data from DynamoDB to the analytics team.

Which solution will meet these requirements?

- A.** Export the required data from the DynamoDB table to an Amazon S3 bucket as multiple JSON files. Provide the analytics team with the necessary IAM permissions to access the S3 bucket.
- B.** Allow public access to the DynamoDB table. Create an IAM user that has permission to access DynamoDB. Share the IAM user with the analytics team.
- C.** Allow public access to the DynamoDB table. Create an IAM user that has read-only permission for DynamoDB. Share the IAM user with the analytics team.
- D.** Create a cross-account IAM role with a permissions policy that allows access to the DynamoDB table. Configure the role trust policy to allow the analytics team's AWS account to assume the role.

#### ANSWER: D

##### Explanation:

Create a cross-account IAM role with permissions to read the required DynamoDB table, and configure its trust policy to allow the analytics account to assume it. This provides secure, temporary credentials to the analytics application instead of requiring long-lived shared credentials. The EC2 workload can use an IAM role in the analytics account and AWS Security Token Service to assume the role in the data-owning account. The assumed role's permissions policy can be scoped to only the required DynamoDB actions, such as `dynamodb:Query`, `dynamodb:Scan`, or `dynamodb:GetItem`, and only the specific table and indexes that the analytics workload needs. This supports regularly accessing the current table data, including frequent updates, without copying sensitive records to a separate location. It also supports least privilege, central auditing through AWS CloudTrail, and simple revocation or adjustment of access by changing the role or its policies. AWS recommends IAM roles for delegating cross-account access because roles use temporary security credentials and establish an explicit trust relationship between the resource-owning and consuming accounts. See [AWS IAM cross-account roles documentation](#) and [Amazon DynamoDB IAM security documentation](#).

#### QUESTION NO: 17

A company is deploying a new application on Amazon EC2 instances. The application writes data to Amazon Elastic Block Store (Amazon EBS) volumes. The company needs to ensure that all data that is written to the EBS volumes is encrypted at rest.

Which solution will meet this requirement?

- A.** Create an IAM role that specifies EBS encryption. Attach the role to the EC2 instances.
- B.** Create the EBS volumes as encrypted volumes. Attach the EBS volumes to the EC2 instances.

C. Create an EC2 instance tag that has a key of Encrypt and a value of True. Tag all instances that require encryption at the ESS level.

D. Create an AWS Key Management Service (AWS KMS) key policy that enforces EBS encryption in the account. Ensure that the key policy is active.

**ANSWER: B**

**Explanation:**

Creating the EBS volumes as encrypted volumes and attaching them to the EC2 instances ensures that data written by the application is encrypted at rest. Amazon EBS encryption is performed transparently: data at rest on the volume, data moving between the EC2 instance and the volume, and EBS snapshots created from the encrypted volume are protected using AWS Key Management Service (AWS KMS) encryption keys. The application and operating system do not need to implement encryption logic or make changes to read and write data on the volume. When creating an encrypted volume, the company can use the default AWS managed key for EBS or select an appropriate customer managed KMS key when it needs control over key administration, access permissions, rotation, or auditing. For a consistent account-wide default, the company can also enable EBS encryption by default; however, directly creating the application volumes as encrypted volumes definitively meets the stated requirement for these instances. Encryption must be selected when the volume is created because an existing unencrypted EBS volume cannot be encrypted in place. To migrate existing data, create an encrypted snapshot and restore it to a new encrypted volume. See the [Amazon EBS encryption documentation](#) and [EBS encryption by default documentation](#).

**QUESTION NO: 18**

A company has a single AWS account that contains resources belonging to several teams. The company needs to identify the costs associated with each team. The company wants to use a tag named CostCenter to identify resources that belong to each team.

A. Tag all resources that belong to each team with the user-defined CostCenter tag.

B. Create a tag for each team, and set the value to CostCenter.

C. Activate the CostCenter tag to track cost allocation.

D. Configure AWS Billing and Cost Management to send monthly invoices to the company through email messages.

E. Set up consolidated billing in the existing AWS account.

**ANSWER: A C**

**Explanation:**

Tag all applicable resources with the user-defined `CostCenter` tag, using a distinct tag value for each team, such as `Engineering`, `Marketing`, or an internal cost-center identifier. AWS records tag values on supported resources and can use them to categorize the associated usage and costs. Consistent tagging is essential: resources must receive the same tag key, `CostCenter`, while the value identifies the responsible team.

The `CostCenter` user-defined cost allocation tag must then be activated in the AWS Billing and Cost Management console. Activation makes the tag available as a cost allocation dimension in AWS cost-management data, including Cost Explorer and cost allocation reports. After activation, AWS can take up to 24 hours for the tag to appear in cost-management products, and cost allocation tag data applies going forward rather than retroactively. These two steps provide the required team-level allocation and reporting within the existing account. See the [AWS documentation for cost allocation tags](#) and [activating user-defined cost allocation tags](#).

**QUESTION NO: 19**

A company has two applications: a sender application that sends messages with payloads to be processed and a processing application intended to receive the messages with payloads. The company wants to implement an AWS service to handle messages between the two applications. The sender application can send about 1,000 messages each hour. The messages

may take up to 2 days to be processed. If the messages fail to process, they must be retained so that they do not impact the processing of any remaining messages.

Which solution meets these requirements and is the MOST operationally efficient?

- A.** Set up an Amazon EC2 instance running a Redis database. Configure both applications to use the instance. Store, process, and delete the messages, respectively.
- B.** Use an Amazon Kinesis data stream to receive the messages from the sender application. Integrate the processing application with the Kinesis Client Library (KCL).
- C.** Integrate the sender and processor applications with an Amazon Simple Queue Service (Amazon SQS) queue. Configure a dead-letter queue to collect the messages that failed to process.
- D.** Subscribe the processing application to an Amazon Simple Notification Service (Amazon SNS) topic to receive notifications to process. Integrate the sender application to write to the SNS topic.

**ANSWER: C**

**Explanation:**

Integrate the sender and processor applications with an Amazon Simple Queue Service (Amazon SQS) queue. Configure a dead-letter queue to collect the messages that failed to process. Amazon SQS is a fully managed message-queuing service designed to decouple producers from consumers. The sender can enqueue messages independently of the processing application's availability or processing rate, while the processor polls and handles messages at its own pace. This eliminates infrastructure management and scales easily for the stated message volume.

An SQS queue supports a message retention period of up to 14 days, so messages can remain available while processing takes as long as two days. The processing application can use a visibility timeout appropriate to its processing behavior and delete a message only after successful completion. A redrive policy can send a message to a dead-letter queue after it exceeds the configured maximum number of receive attempts. The dead-letter queue isolates repeatedly failing messages from the source queue, allowing healthy messages to continue processing without blockage while preserving failed messages for diagnosis and later remediation. See the [Amazon SQS dead-letter queue documentation](#) and the [Amazon SQS queue parameter documentation](#).

## QUESTION NO: 20

A company stores petabytes of historical medical information on premises. The company has a process to manage encryption of the data to comply with regulations. The company needs a cloud-based solution for data backup, recovery, and archiving. The company must retain control over the encryption key material. Which combination of solutions will meet these requirements? (Select TWO.)

- A.** Create an AWS Key Management Service (AWS KMS) key without key material. Import the company 's key material into the KMS key.
- B.** Create an AWS Key Management Service (AWS KMS) encryption key that contains key material generated by AWS KMS.
- C.** Store the data in Amazon S3 Standard-Infrequent Access (S3 Standard-IA) storage. Use S3 Bucket Keys with AWS Key Management Service (AWS KMS) keys.
- D.** Store the data in an Amazon S3 Glacier storage class. Use server-side encryption with customer-provided keys (SSE-C).
- E.** Store the data in AWS Snowball devices. Use server-side encryption with AWS KMS keys (SSE-KMS).
- F.** Encrypt the data client-side with company-managed encryption keys before storing it in an Amazon S3 Glacier storage class.

**ANSWER: A F**

**Explanation:**

Creating an AWS KMS key without key material and importing the company's key material enables the company to use its own existing key material with AWS KMS. The company retains the original key material and controls its lifecycle, including when to delete the imported material from AWS KMS. This supports a migration to AWS while allowing the organization to maintain its established key-management process. AWS documents the requirements and lifecycle behavior for imported key material at [Import key material into AWS KMS](#).

Storing the data in an Amazon S3 Glacier storage class after encrypting it client-side provides the required low-cost cloud archive, backup, and recovery destination while retaining encryption-key control entirely with the company. With client-side encryption, data is encrypted before it is sent to Amazon S3, and the organization generates, manages, and retains the encryption keys outside AWS. The encrypted objects can then be archived using an S3 Glacier storage class without AWS needing access to the plaintext encryption keys. Amazon S3 describes this model at [Protecting data using client-side encryption](#).

#### QUESTION NO: 21

A company is migrating an on-premises Oracle database to Amazon Aurora PostgreSQL. The company wants to minimize downtime during the migration. The source database must remain available while initial data is loaded and while ongoing changes are replicated to the target database.

Which combination of actions should a solutions architect take to meet these requirements? (Choose two.)

- A. Use the AWS Schema Conversion Tool (AWS SCT) to convert the Oracle schema and code objects for Aurora PostgreSQL.
- B. Create an AWS Database Migration Service (AWS DMS) replication instance and an ongoing replication task.
- C. Create an Amazon RDS for Oracle read replica and promote the read replica as Aurora PostgreSQL.
- D. Create an Amazon S3 bucket and export Oracle database backups to the bucket every hour.
- E. Use Amazon ElastiCache for Redis to replicate the Oracle database changes to Aurora PostgreSQL.

#### ANSWER: A B

##### Explanation:

Use the AWS Schema Conversion Tool (AWS SCT) to convert the Oracle database schema and code objects to a format compatible with Aurora PostgreSQL. Oracle and PostgreSQL are heterogeneous database engines, so schema objects and Oracle-specific code commonly require assessment and conversion before the migration.

Create an AWS Database Migration Service (AWS DMS) replication instance and an ongoing replication task. AWS DMS can perform an initial full load and then use change data capture (CDC) to replicate ongoing source changes to Aurora PostgreSQL. This keeps the target synchronized while the Oracle source remains available until the company is ready to cut over. See the [AWS Schema Conversion Tool documentation](#) and the [AWS DMS change data capture documentation](#).

#### QUESTION NO: 22

A company runs a high-traffic web application that has a three-tier architecture consisting of a web layer, an application layer, and a database layer. The web layer and application layer run on Amazon EC2 instances behind an Application Load Balancer (ALB). The application layer is stateless and supports automatic scaling. The database layer uses Amazon RDS for MySQL in a Multi-AZ configuration and relies on a relational architecture.

The company is preparing for a large marketing event that is expected to drive a sharp increase in read traffic. The company must ensure that the application remains highly available and responsive under load. The company wants to scale the application's architecture components but does not want to modify the application.

Which combination of solutions will meet these requirements? (Select THREE.)

- A. Deploy an Amazon CloudFront distribution. Specify the web layer as the origin.
- B. Enable automatic scaling for EC2 instances in the application layer.

- C. Migrate the database to Amazon Aurora. Configure Aurora Auto Scaling and Aurora Replicas.
- D. Set up an Amazon ElastiCache (Redis OSS) cluster in front of the database.
- E. Replace the ALB with a Network Load Balancer (NLB).
- F. Migrate the database to an Amazon DynamoDB table.

**ANSWER: A B C**

**Explanation:**

Deploy an Amazon CloudFront distribution. Specify the web layer as the origin. CloudFront caches cacheable content at edge locations close to viewers, reducing requests that must be served by the web tier and improving response times during a traffic surge. The existing Application Load Balancer can be used as the CloudFront origin endpoint for the web tier, retaining the current highly available load-balanced design.

Enable automatic scaling for EC2 instances in the application layer. Because this tier is stateless, instances can be added or removed safely in response to demand. An EC2 Auto Scaling group behind the Application Load Balancer provides elastic capacity while the load balancer distributes incoming requests across healthy instances. This directly supports responsiveness and availability without requiring application changes.

Migrate the database to Amazon Aurora. Configure Aurora Auto Scaling and Aurora Replicas. Aurora is compatible with MySQL, enabling the company to retain its relational model while improving read scalability. Aurora Replicas provide additional read endpoints to offload read traffic from the writer, and Aurora Auto Scaling adjusts the number of replicas according to demand. Aurora also provides storage replication across multiple Availability Zones, supporting the availability requirement. See the [Amazon CloudFront Developer Guide](#) and [Amazon Aurora replication documentation](#).

**QUESTION NO: 23**

A company regularly uploads GB-sized files to Amazon S3. After the company uploads the files, the company uses a fleet of Amazon EC2 Spot Instances to transcode the file format. The company needs to scale throughput when the company uploads data from the on-premises data center to Amazon S3 and when the company downloads data from Amazon S3 to the EC2 instances.

Which solutions will meet these requirements? (Select TWO.)

- A. Use the S3 bucket access point instead of accessing the S3 bucket directly.
- B. Upload the files into multiple S3 buckets.
- C. Use S3 multipart uploads.
- D. Fetch multiple byte-ranges of an object in parallel.
- E. Add a random prefix to each object when uploading the files.

**ANSWER: C D**

**Explanation:**

**Use S3 multipart uploads** is correct because multipart upload divides a large object into independently uploaded parts. The on-premises uploader can transfer multiple parts concurrently, which increases aggregate upload throughput and makes large transfers more resilient: if a part fails, only that part needs to be retried. Amazon S3 supports multipart upload for objects and recommends considering it for objects of 100 MB or larger; it is particularly appropriate for the GB-sized files in this scenario.

**Fetch multiple byte-ranges of an object in parallel.** is correct because the transcoding fleet can issue concurrent GET requests for distinct byte ranges of the same S3 object. S3 supports byte-range fetches through the HTTP `Range` header. Parallelizing those range requests increases the effective download throughput available to an EC2 instance or processing worker, enabling the fleet to retrieve large source files more efficiently before or during transcoding. Together, multipart uploads and parallel byte-range retrieval directly address scalable transfer throughput in each required direction without requiring changes to the object namespace or additional buckets.

### QUESTION NO: 24

A company is developing an ecommerce application that will consist of a load-balanced front end, a container-based application, and a relational database. A solutions architect needs to create a highly available solution that operates with as little manual intervention as possible.

Which solutions meet these requirements? Select TWO.

- A. Create an Amazon RDS DB instance in Multi-AZ mode.
- B. Create an Amazon RDS DB instance and one or more replicas in another Availability Zone.
- C. Create an Amazon EC2 instance-based Docker cluster to handle the dynamic application load.
- D. Create an Amazon ECS cluster with a Fargate launch type to handle the dynamic application load.
- E. Create an Amazon ECS cluster with an Amazon EC2 launch type to handle the dynamic application load.

### ANSWER: A D

#### Explanation:

Creating an Amazon RDS DB instance in Multi-AZ mode provides a managed high-availability deployment for the relational database tier. RDS synchronously replicates data to a standby DB instance in a different Availability Zone. If the primary DB instance, its Availability Zone, or certain underlying infrastructure experiences a failure, Amazon RDS performs automatic failover to the standby. The application can reconnect using the DB instance endpoint, minimizing administrative work during a failure event. This design directly addresses availability while retaining the operational benefits of a managed database service.

Creating an Amazon ECS cluster with a Fargate launch type to handle the dynamic application load minimizes management of the container compute layer. With Fargate, AWS provisions and manages the servers that run container tasks, so the company does not need to select, patch, maintain, or scale an EC2 worker fleet. ECS services can maintain the desired number of healthy tasks and replace failed tasks automatically. When tasks are deployed across multiple Availability Zones and connected to the load-balanced front end, this provides a resilient application tier with low operational overhead. See the [Amazon RDS Multi-AZ documentation](#) and the [AWS Fargate documentation](#).

### QUESTION NO: 25

A company has thousands of edge devices that collectively generate 1 TB of status alerts each day. Each alert is approximately 2 KB in size. A solutions architect needs to implement a solution to ingest and store the alerts for future analysis.

The company wants a highly available solution. However, the company needs to minimize costs and does not want to manage additional infrastructure. Additionally, the company wants to keep 14 days of data available for immediate analysis and archive any data older than 14 days.

What is the MOST operationally efficient solution that meets these requirements?

- A. Create an Amazon Kinesis Data Firehose delivery stream to ingest the alerts. Configure the Kinesis Data Firehose stream to deliver the alerts to an Amazon S3 bucket. Set up an S3 Lifecycle configuration to transition data to Amazon S3 Glacier after 14 days.
- B. Launch Amazon EC2 instances across two Availability Zones and place them behind an Elastic Load Balancer to ingest the alerts. Create a script on the EC2 instances that will store the alerts in an Amazon S3 bucket. Set up an S3 Lifecycle configuration to transition data to Amazon S3 Glacier after 14 days.
- C. Create an Amazon Kinesis Data Firehose delivery stream to ingest the alerts. Configure the Kinesis Data Firehose stream to deliver the alerts to an Amazon Elasticsearch Service (Amazon ES) cluster. Set up the Amazon ES cluster to take manual snapshots every day and delete data from the cluster that is older than 14 days.

**D.** Create an Amazon Simple Queue Service (Amazon SQS) standard queue to ingest the alerts and set the message retention period to 14 days. Configure consumers to poll the SQS queue, check the age of the message, and analyze the message data as needed. If the message is 14 days old, the consumer should copy the message to an Amazon S3 bucket and delete the message from the SQS queue.

**ANSWER: A**

**Explanation:**

Creating an Amazon Kinesis Data Firehose delivery stream to ingest the alerts, delivering them to Amazon S3, and using an S3 Lifecycle configuration to transition objects to Amazon S3 Glacier after 14 days is the most operationally efficient solution. Kinesis Data Firehose is a fully managed ingestion and delivery service that automatically scales to accommodate the high aggregate volume of small device alerts. It removes the need to provision, patch, scale, or operate servers while providing resilient managed delivery to S3.

Amazon S3 provides highly durable, highly available storage for the incoming alert data and supports direct analysis during the first 14 days. For example, the retained S3 data can be queried with services such as Amazon Athena without moving it elsewhere. An S3 Lifecycle rule can then automatically transition objects after 14 days to an Amazon S3 Glacier storage class, meeting the archival requirement without custom jobs or ongoing operational intervention. This combines serverless ingestion, low-cost durable storage, and automated age-based data management. See the [Amazon Data Firehose documentation](#) and the [Amazon S3 lifecycle management documentation](#).

**QUESTION NO: 26**

A company plans to use AWS to run high-performance computing (HPC) workloads and analytics workloads. The company will run HPC workloads on Amazon EC2 instances. The workloads require a high-performance file system that can scale to millions of input/output operations per second (IOPS). Which combination of steps will meet these requirements? (Select TWO.)

- A.** Use Amazon Elastic File System (Amazon EFS) as a high-performance file system.
- B.** Use Amazon FSx for Lustre as a high-performance file system.
- C.** Create an Auto Scaling group of Amazon EC2 instances. Use Reserved Instances. Configure a spread placement group. Use AWS Batch to run the analytics workloads.
- D.** Use Mountpoint for Amazon S3 as a high-performance file system.
- E.** Create an Auto Scaling group of Amazon EC2 instances. Use a mix of On-Demand Instances, Reserved Instances, and Spot Instances. Configure a cluster placement group. Use Amazon EMR to run the analytics workloads.

**ANSWER: B E**

**Explanation:**

Using Amazon FSx for Lustre as a high-performance file system provides the parallel, high-throughput shared file storage required for EC2-based HPC workloads. FSx for Lustre is designed for compute-intensive use cases such as HPC, machine learning, video processing, and financial modeling. It integrates with Amazon S3 for input and output data and can scale performance independently with its file-system configuration. AWS documents that FSx for Lustre file systems can deliver millions of IOPS and hundreds of GBps of throughput, allowing many compute instances to access a common data set concurrently. See the [Amazon FSx for Lustre User Guide](#).

Creating an Auto Scaling group of Amazon EC2 instances, using a mix of On-Demand Instances, Reserved Instances, and Spot Instances, configuring a cluster placement group, and using Amazon EMR for analytics addresses the compute, networking, cost, and analytics portions of the requirement. A cluster placement group places instances close together within a single Availability Zone, supporting the low-latency, high-bandwidth networking commonly needed by tightly coupled HPC jobs. Auto Scaling supports capacity management, while the instance purchasing mix can optimize cost for workloads with differing interruption and capacity needs. Amazon EMR is a managed platform for large-scale distributed analytics. AWS placement-group guidance is available in the [Amazon EC2 User Guide](#).

## QUESTION NO: 27

A company needs a solution to prevent photos with unwanted content from being uploaded to the company's web application. The solution must not involve training a machine learning (ML) model.

Which solution will meet these requirements?

- A.** Create and deploy a model by using Amazon SageMaker Autopilot. Create a real-time endpoint that the web application invokes when new photos are uploaded.
- B.** Create an AWS Lambda function that uses Amazon Rekognition to detect unwanted content. Create a Lambda function URL that the web application invokes when new photos are uploaded.
- C.** Create an Amazon CloudFront function that uses Amazon Comprehend to detect unwanted content. Associate the function with the web application.
- D.** Create an AWS Lambda function that uses Amazon Rekognition Video to detect unwanted content. Create a Lambda function URL that the web application invokes when new photos are uploaded.

## ANSWER: B

### Explanation:

Create an AWS Lambda function that uses Amazon Rekognition to detect unwanted content. Create a Lambda function URL that the web application invokes when new photos are uploaded. This solution uses Amazon Rekognition's managed image-moderation capability, specifically the `DetectModerationLabels` API, to evaluate still images for categories of potentially inappropriate, unsafe, or otherwise unwanted visual content. Rekognition supplies and operates the underlying pretrained ML models, so the company does not need to collect labeled images, build a model, train it, tune it, or manage inference infrastructure.

The Lambda function can receive the upload request, or an image reference such as an Amazon S3 object key, invoke Rekognition, and compare returned moderation labels and confidence scores with the company's policy. It can then allow the upload to proceed, quarantine the image, or return a rejection response. A Lambda function URL gives the web application an HTTPS endpoint for invoking that logic while retaining a serverless and automatically scaling implementation. This pattern directly addresses photo moderation with low operational overhead and no customer-managed ML training. See the [Amazon Rekognition image moderation documentation](#) and the [AWS Lambda function URLs documentation](#).

## QUESTION NO: 28

A healthcare company is developing an AWS Lambda function that publishes notifications to an encrypted Amazon Simple Notification Service (Amazon SNS) topic. The notifications contain protected health information (PHI).

The SNS topic uses AWS Key Management Service (AWS KMS) customer-managed keys for encryption. The company must ensure that the application has the necessary permissions to publish messages securely to the SNS topic.

Which combination of steps will meet these requirements? (Select THREE.)

- A.** Create a resource policy for the SNS topic that allows the Lambda function to publish messages to the topic.
- B.** Use server-side encryption with AWS KMS keys (SSE-KMS) for the SNS topic instead of customer-managed keys.
- C.** Create a resource policy for the encryption key that the SNS topic uses that has the necessary AWS KMS permissions.
- D.** Specify the Lambda function's Amazon Resource Name (ARN) in the SNS topic's resource policy.
- E.** Associate an Amazon API Gateway HTTP API with the SNS topic to control access to the topic by using API Gateway resource policies.
- F.** Configure a Lambda execution role that has the necessary IAM permissions to use a customer-managed key in AWS KMS.

## ANSWER: A C F

### Explanation:

Creating a resource policy for the SNS topic that allows the Lambda function to publish messages to the topic provides authorization at the topic resource boundary for publishing. In practice, the policy principal should identify the Lambda execution-role principal that invokes the SNS Publish API. This is particularly useful when a topic policy is required to explicitly control who can publish to a sensitive PHI-bearing topic.

Creating a resource policy for the encryption key that the SNS topic uses that has the necessary AWS KMS permissions enables use of the customer-managed KMS key for SNS server-side encryption. The key policy is the primary authorization control for a customer-managed key and must permit the required KMS cryptographic operations. AWS documents the KMS permissions required when using server-side encryption for SNS topics in its [Amazon SNS server-side encryption documentation](#).

Configuring a Lambda execution role that has the necessary IAM permissions to use a customer-managed key in AWS KMS gives the application principal the required identity-based authorization for the KMS operations associated with publishing to an encrypted topic. The execution role is also the correct AWS identity through which Lambda obtains permissions at runtime. AWS explains that effective access to a KMS key depends on both the key policy and IAM authorization in the [AWS KMS key policies documentation](#).

### QUESTION NO: 29

A company has a single AWS account. The company runs workloads on Amazon EC2 instances in multiple VPCs in one AWS Region. The company also runs workloads in an on-premises data center that connects to the company's AWS account by using AWS Direct Connect.

The company needs all EC2 instances in the VPCs to resolve DNS queries for the internal.example.com domain to the authoritative DNS server that is located in the on-premises data center. The solution must use private communication between the VPCs and the on-premises network. All route tables, network ACLs, and security groups are configured correctly between AWS and the on-premises data center.

Which combination of actions will meet these requirements? (Select THREE.)

- A. Create an Amazon Route 53 inbound endpoint in all the workload VPCs.
- B. Create an Amazon Route 53 outbound endpoint in one of the workload VPCs.
- C. Create an Amazon Route 53 Resolver rule with the Forward type configured to forward queries for internal.example.com to the on-premises DNS server.
- D. Create an Amazon Route 53 Resolver rule with the System type configured to forward queries for internal.example.com to the on-premises DNS server.
- E. Associate the Amazon Route 53 Resolver rule with all the workload VPCs.
- F. Associate the Amazon Route 53 Resolver rule with the workload VPC with the new Route 53 endpoint.

### ANSWER: B C E

### Explanation:

Creating an Amazon Route 53 outbound endpoint in one of the workload VPCs provides the managed, highly available path through which Route 53 Resolver can send DNS requests from AWS to DNS resolvers on the connected on-premises network. The endpoint uses elastic network interfaces in selected subnets, and DNS queries travel over the existing private Direct Connect connectivity rather than through the public internet. AWS recommends deploying the endpoint across multiple Availability Zones for resilience.

Creating an Amazon Route 53 Resolver rule with the Forward type configured to forward queries for internal.example.com to the on-premises DNS server defines the conditional forwarding behavior. Requests for names within that domain are sent to the configured IP addresses of the authoritative or forwarding DNS servers in the data center through the outbound endpoint.

Associating the Amazon Route 53 Resolver rule with all the workload VPCs makes that conditional forwarding rule available to every VPC that contains EC2 workloads. Resolver rules are regional resources and can be associated with multiple VPCs in the same Region, allowing a single outbound endpoint to support the required hybrid DNS resolution pattern. See the [Amazon Route 53 Resolver outbound query documentation](#) and [Resolver rule management documentation](#).

### QUESTION NO: 30

A company uses an on-premises network-attached storage (NAS) system to provide file shares to its high performance computing (HPC) workloads. The company wants to migrate its latency-sensitive HPC workloads and its storage to the AWS Cloud. The company must be able to provide NFS and SMB multi-protocol access from the file system.

Which solution will meet these requirements with the LEAST latency? (Select TWO.)

- A. Deploy compute optimized EC2 instances into a cluster placement group.
- B. Deploy compute optimized EC2 instances into a partition placement group.
- C. Attach the EC2 instances to an Amazon FSx for Lustre file system.
- D. Attach the EC2 instances to an Amazon FSx for OpenZFS file system.
- E. Attach the EC2 instances to an Amazon FSx for NetApp ONTAP file system.

**ANSWER: A E**

#### Explanation:

Deploying compute optimized EC2 instances into a cluster placement group helps meet the low-latency requirement for the HPC compute tier. A cluster placement group packs instances close together within a single Availability Zone, enabling very low network latency and high packet-per-second performance between participating instances. This placement strategy is specifically suited to tightly coupled, network-intensive HPC applications. Compute optimized instances provide a high ratio of vCPUs to memory and are appropriate for workloads that need substantial CPU capacity, such as scientific modeling, simulation, and other compute-bound HPC processing.

Attaching the EC2 instances to an Amazon FSx for NetApp ONTAP file system satisfies the file-sharing protocol requirement. FSx for ONTAP provides fully managed NetApp ONTAP file systems with multiprotocol access, including both NFS and SMB. This allows Linux- and Windows-based HPC clients to access the same storage using their required native protocols, while retaining ONTAP data-management capabilities. Using compute resources colocated in an Availability Zone with managed, high-performance file storage provides the appropriate low-latency AWS architecture for the stated workload. See the [Amazon EC2 cluster placement group documentation](#) and the [Amazon FSx for NetApp ONTAP multiprotocol access documentation](#).

### QUESTION NO: 31

A company uses high concurrency AWS Lambda functions to process a constantly increasing number of messages in a message queue during marketing events. The Lambda functions use CPU intensive code to process the messages. The company wants to reduce the compute costs and to maintain service latency for its customers.

Which solution will meet these requirements?

- A. Configure reserved concurrency for the Lambda functions. Decrease the memory allocated to the Lambda functions.
- B. Configure reserved concurrency for the Lambda functions. Increase the memory according to AWS Compute Optimizer recommendations.
- C. Configure provisioned concurrency for the Lambda functions. Decrease the memory allocated to the Lambda functions.
- D. Configure provisioned concurrency for the Lambda functions. Increase the memory according to AWS Compute Optimizer recommendations.

**ANSWER: D**

#### Explanation:

Configure provisioned concurrency for the Lambda functions. Increase the memory according to AWS Compute Optimizer recommendations. Provisioned concurrency keeps a specified number of Lambda execution environments initialized and ready to serve requests. This minimizes cold-start latency during marketing-event traffic spikes, helping the message-

processing service remain responsive as concurrency increases. It is particularly useful when consistent low latency is required for workloads with variable or bursty invocation rates.

For Lambda, CPU capacity scales proportionally with the configured memory setting. CPU-intensive processing can therefore complete more quickly when memory is appropriately increased, which can reduce execution duration and potentially lower total cost even if the per-millisecond rate is higher. AWS Compute Optimizer analyzes historical Lambda utilization metrics and provides memory-size recommendations that balance performance and cost for the observed workload. Applying its recommendation helps avoid selecting memory solely by guesswork and supports cost-efficient throughput as message volume grows.

Together, pre-initialized execution capacity protects customer-facing latency, while a data-driven memory configuration supplies suitable CPU resources and can reduce duration-based Lambda charges. See [AWS Lambda provisioned concurrency documentation](#) and [AWS Compute Optimizer documentation](#).

### QUESTION NO: 32

A company uses Amazon EC2 instances behind an Application Load Balancer ALB to serve content to users. The company uses Amazon EBS volumes to store data.

The company needs to encrypt data in transit and at rest.

Which combination of services will meet these requirements? Select TWO.

- A. Amazon GuardDuty
- B. AWS Shield
- C. AWS Certificate Manager ACM
- D. AWS Secrets Manager
- E. AWS KMS

### ANSWER: C E

#### Explanation:

AWS Certificate Manager ACM provides and manages public and private TLS/SSL certificates. A certificate issued or imported into ACM can be attached to an HTTPS listener on an Application Load Balancer, allowing the load balancer to encrypt traffic in transit between users and the ALB. ACM also handles managed certificate renewal for eligible public certificates, reducing the operational burden of maintaining certificates.

AWS KMS provides the customer managed or AWS managed keys that Amazon EBS uses for encryption at rest. When EBS encryption is enabled, EBS encrypts volume data, disk I/O, snapshots created from encrypted volumes, and volumes created from those snapshots using the selected KMS key. This meets the requirement to protect the data stored on the EC2 instances' EBS volumes while retaining key-policy and audit controls through KMS and AWS CloudTrail.

Together, AWS Certificate Manager ACM and AWS KMS address the distinct encryption boundaries in the workload: TLS for client connections to the Application Load Balancer and KMS-backed encryption for persistent block storage. See the [Application Load Balancer HTTPS listener documentation](#) and the [Amazon EBS encryption documentation](#).

### QUESTION NO: 33

A company is building a new web-based customer relationship management application. The application will use several Amazon EC2 instances that are backed by Amazon EBS volumes behind an Application Load Balancer (ALB). The application will also use an Amazon Aurora database. All data for the application must be encrypted at rest and in transit.

Which solution will meet these requirements?

- A. Use AWS KMS certificates on the ALB to encrypt data in transit. Use AWS Certificate Manager (ACM) to encrypt the EBS volumes and Aurora database storage at rest.

**B.** Use the AWS root account to log in to the AWS Management Console. Upload the company ' s encryption certificates. While in the root account, select the option to turn on encryption for all data at rest and in transit for the account.

**C.** Use AWS KMS to encrypt the EBS volumes and Aurora database storage at rest. Attach an AWS Certificate Manager (ACM) certificate to the ALB to encrypt data in transit.

**D.** Use BitLocker to encrypt all data at rest. Import the company ' s TLS certificate keys to AWS KMS. Attach the KMS keys to the ALB to encrypt data in transit.

**ANSWER: C**

**Explanation:**

Use AWS KMS to encrypt the EBS volumes and Aurora database storage at rest. Attach an AWS Certificate Manager (ACM) certificate to the ALB to encrypt data in transit. This solution maps AWS-managed encryption capabilities to the appropriate layer of the application. Amazon EBS encryption uses AWS Key Management Service (AWS KMS) keys to protect volume data at rest, including data on the volume and associated snapshots. Amazon Aurora also supports encryption at rest using KMS keys, protecting the underlying database storage, backups, snapshots, and automated backups.

For web traffic in transit, an Application Load Balancer HTTPS listener uses an SSL/TLS certificate. AWS Certificate Manager can issue, manage, and deploy public or private certificates for supported AWS services, including Elastic Load Balancing. Associating an ACM certificate with the ALB HTTPS listener enables TLS encryption for client connections to the application. Where traffic must remain encrypted beyond the load balancer, the application architecture can use HTTPS target connections and TLS database connections as appropriate. This approach uses managed AWS services for centralized key control and certificate lifecycle management while satisfying the stated encryption-at-rest and encryption-in-transit design requirements. See the [Amazon EBS encryption documentation](#) and the [Application Load Balancer HTTPS listener documentation](#).

**QUESTION NO: 34**

A company has a data ingestion workflow that consists the following:

- An Amazon Simple Notification Service (Amazon SNS) topic for notifications about new data deliveries
- An AWS Lambda function to process the data and record metadata

The company observes that the ingestion workflow fails occasionally because of network connectivity issues. When such a failure occurs, the Lambda function does not ingest the corresponding data unless the company manually reruns the job.

Which combination of actions should a solutions architect take to ensure that the Lambda function ingests all data in the future? (Select TWO.)

- A.** Configure the Lambda function In multiple Availability Zones.
- B.** Create an Amazon Simple Queue Service (Amazon SQS) queue, and subscribe It to me SNS topic.
- C.** Increase the CPU and memory that are allocated to the Lambda function.
- D.** Increase provisioned throughput for the Lambda function.
- E.** Modify the Lambda function to read from an Amazon Simple Queue Service (Amazon SQS) queue

**ANSWER: B E**

**Explanation:**

Creating an Amazon SQS queue and subscribing it to the Amazon SNS topic provides a durable buffer between notification delivery and processing. Each SNS notification is persisted in the queue until it is successfully consumed or reaches the queue's configured retention period. This decouples the publisher from the ingestion processor, so a transient networking failure or temporary Lambda processing issue does not cause a notification to be lost.

Modifying the Lambda function to read from the Amazon SQS queue makes SQS the Lambda event source. Lambda polls the queue and invokes the function with queued messages. If processing fails, Lambda and SQS provide retry behavior: the

message remains available for later processing after the visibility timeout expires. The queue can also be configured with an appropriate visibility timeout, retention period, and a dead-letter queue for messages that repeatedly fail. Together, these actions support reliable, asynchronous ingestion and remove the need for manual reruns after temporary connectivity failures. Because SQS delivery is at least once, the ingestion function should also be designed to handle duplicate messages safely, such as by making metadata updates idempotent. See the AWS documentation for [using Lambda with Amazon SQS](#) and [subscribing an SQS queue to an SNS topic](#).

### QUESTION NO: 35

A media company has a multi-account AWS environment in the us-east-1 Region. The company has an Amazon Simple Notification Service (Amazon SNS) topic in a production account that publishes performance metrics. The company has an AWS Lambda function in an administrator account to process and analyze log data.

The Lambda function that is in the administrator account must be invoked by messages from the SNS topic that is in the production account when significant metrics  $tM^*$  reported.

Which combination of steps will meet these requirements? (Select TWO.)

- A. Create a resource-based policy for the Lambda function that allows Amazon SNS to invoke the function, restricted to the SNS topic ARN in the production account.
- B. Create an SNS topic resource policy that allows the administrator account to subscribe its Lambda function to the topic.
- C. Use an Amazon EventBridge rule in the production account to capture the SNS topic notifications. Configure the EventBridge rule to forward notifications to the Lambda function that is in the administrator account.
- D. Store performance metrics in an Amazon S3 bucket in the production account. Use Amazon Athena to analyze the metrics from the administrator account.

### ANSWER: A B

#### Explanation:

Cross-account Amazon SNS-to-Lambda delivery requires authorization at both the SNS topic and the Lambda function. Create a resource-based policy for the Lambda function that allows the Amazon SNS service principal to invoke the function, scoped to the ARN of the production account's topic. This grants the service permission to perform `lambda:InvokeFunction` only when the invocation originates from the intended SNS topic, which is the required Lambda-side authorization for an SNS subscription.

The SNS topic also needs a topic resource policy that permits the administrator account to create a subscription for its Lambda function. With that cross-account `sns:Subscribe` permission in place, the administrator account can subscribe the function endpoint to the production topic. SNS then delivers each published significant-metric notification to the subscribed Lambda function, and Lambda invokes the configured function asynchronously. Both the topic policy and the Lambda resource-based permission should use least privilege by naming the specific account and topic ARN rather than granting broad access.

AWS documents the SNS subscription and Lambda invocation permission model in [Invoking a Lambda function using Amazon SNS notifications](#) and the required Lambda permissions in [Using Lambda with Amazon SNS](#).

### QUESTION NO: 36

A company requires centralized auditing for all AWS accounts and compliance monitoring against AWS Foundational Security Best Practices (FSBP) with minimal operational overhead.

Which solution will meet these requirements?

- A. Deploy AWS Control Tower in the management account. Enable AWS Security Hub and Account Factory.
- B. Deploy AWS Control Tower in a member account.
- C. Use AWS Managed Services (AMS) with GuardDuty.

**D. Use AWS Managed Services (AMS) with Security Hub.**

**ANSWER: A**

**Explanation:**

Deploy AWS Control Tower in the management account. Enable AWS Security Hub and Account Factory. AWS Control Tower is designed to establish and govern a multi-account AWS environment through AWS Organizations from the organization's management account. During landing-zone setup, Control Tower establishes centralized logging capabilities, including organization-wide AWS CloudTrail logging and AWS Config recording, with logs stored in dedicated Log Archive and Audit accounts. This provides the required centralized auditing foundation without requiring the company to independently design, deploy, and maintain those components across every account.

AWS Security Hub provides the compliance-monitoring capability. When the AWS Foundational Security Best Practices standard is enabled, Security Hub automatically runs its included controls and produces findings and compliance status based on AWS Config configuration data and supported AWS service integrations. Security Hub can aggregate findings across organization accounts and Regions, allowing security teams to review the organization's posture centrally. Account Factory supports consistent, governed provisioning of additional member accounts into the Control Tower environment. Together, these managed services minimize operational effort while supplying both standardized account governance and continuous FSBP compliance visibility. See the [AWS Control Tower User Guide](#) and the [AWS Security Hub FSBP standard documentation](#).

**QUESTION NO: 37**

A company receives data transfers from a small number of external clients that use SFTP software on an Amazon EC2 instance. The clients use an SFTP client to upload data. The clients use SSH keys for authentication. Every hour, an automated script transfers new uploads to an Amazon S3 bucket for processing.

The company wants to move the transfer process to an AWS managed service and to reduce the time required to start data processing. The company wants to retain the existing user management and SSH key generation process. The solution must not require clients to make significant changes to their existing processes.

Which solution will meet these requirements?

- A.** Reconfigure the script that runs on the EC2 instance to run every 15 minutes. Create an S3 Event Notifications rule for all new object creation events. Set an Amazon Simple Notification Service (Amazon SNS) topic as the destination.
- B.** Create an AWS Transfer Family SFTP server that uses the existing S3 bucket as a target. Use service-managed users to enable authentication.
- C.** Require clients to add the AWS DataSync agent into their local environments. Create an IAM user for each client that has permission to upload data to the target S3 bucket.
- D.** Create an AWS Transfer Family SFTP connector that has permission to access the target S3 bucket for each client. Store credentials in AWS Systems Manager. Create an IAM role to allow the SFTP connector to securely use the credentials.

**ANSWER: B**

**Explanation:**

Create an AWS Transfer Family SFTP server that uses the existing S3 bucket as a target. Use service-managed users to enable authentication. AWS Transfer Family provides a fully managed SFTP endpoint that external clients can use with their existing standard SFTP software. The clients can continue authenticating with their SSH public keys, so the established key-generation workflow and client-side process can remain substantially unchanged. Service-managed users let the company create and manage Transfer Family user identities and associate each user with an SSH public key and an IAM role that controls access to the appropriate Amazon S3 location.

With Amazon S3 configured as the server's storage backend, files are written directly to the target bucket when clients upload them. This removes the EC2-hosted SFTP server and eliminates the hourly script that previously copied files into S3. As a result, downstream processing can begin as soon as the new object is available, such as through Amazon S3 event notifications or event-driven processing services. AWS manages the SFTP server infrastructure, availability, patching, and

scaling, while the company retains control of users, keys, and S3 permissions. See the [AWS Transfer Family SFTP server documentation](#) and the [service-managed users documentation](#).

#### QUESTION NO: 38

A company runs an internet-facing web application on AWS and uses Amazon Route 53 with a public hosted zone.

The company wants to log DNS response codes to support future root cause analysis.

Which solution will meet these requirements?

- A. Use Route 53 to configure query logging.
- B. Use AWS CloudTrail to record all Route 53 queries.
- C. Use Amazon CloudWatch metrics for Route 53.
- D. Use AWS Trusted Advisor for root cause analysis.

#### ANSWER: A

##### Explanation:

Use Route 53 to configure query logging. Route 53 public hosted zone query logging records the DNS queries that Route 53 receives for the hosted zone and sends the logs to an Amazon CloudWatch Logs log group. Each log entry includes key diagnostic fields for a DNS request, including the query name, query type, requester IP address, Route 53 edge location, transport protocol, and the DNS response code. The response-code field provides the needed evidence for investigating outcomes such as successful resolutions, nonexistent domains, and server failures.

Because the logs are retained in CloudWatch Logs, the company can select an appropriate retention period and use CloudWatch Logs Insights to search and correlate historical DNS events during an incident investigation. This makes query logging appropriate for future root cause analysis of an internet-facing application that uses a Route 53 public hosted zone. Route 53 documentation describes the query-log fields, including `response_code`, and the process for associating a public hosted zone with a query logging configuration. See the [Amazon Route 53 query logging documentation](#) and the [Route 53 query log format reference](#).

#### QUESTION NO: 39

A company runs multiple workloads on virtual machines (VMs) in an on-premises data center. The company is expanding rapidly. The on-premises data center is not able to scale fast enough to meet business needs. The company wants to migrate the workloads to AWS.

The migration is time sensitive. The company wants to use a lift-and-shift strategy for non-critical workloads.

Which combination of steps will meet these requirements? (Select THREE.)

- A. Use the AWS Schema Conversion Tool (AWS SCT) to collect data about the VMs.
- B. Use AWS Application Migration Service. Install the AWS Replication Agent on the VMs.
- C. Complete the initial replication of the VMs. Launch test instances to perform acceptance tests on the VMs.
- D. Stop all operations on the VMs. Launch a cutover instance.
- E. Use AWS App2Container (A2C) to collect data about the VMs.
- F. Use AWS Database Migration Service (AWS DMS) to migrate the VMs.

#### ANSWER: B C D

##### Explanation:

**Use AWS Application Migration Service. Install the AWS Replication Agent on the VMs., Complete the initial replication of the VMs. Launch test instances to perform acceptance tests on the VMs., and Stop all operations on the VMs Launch a cutover instance.** form the standard AWS Application Migration Service workflow for a time-sensitive rehost, or lift-and-shift, migration. AWS Application Migration Service is designed to migrate physical, virtual, and cloud-based servers to AWS with minimal disruption. Installing the AWS Replication Agent on each source VM enables continuous block-level replication of the source server disks into the AWS staging area. After initial synchronization is complete, the migration team can launch non-production test instances from replicated data and validate application behavior, networking, and operational acceptance before migration. When the workload is ready to move, stopping source activity prevents new writes from creating divergence during the final synchronization. The team can then launch cutover instances, which creates the target Amazon EC2 instances using the latest replicated server data. This sequence supports rapid migration while retaining a testing checkpoint before production cutover. See the [AWS Application Migration Service User Guide](#) and the [AWS documentation on testing and cutover](#).

#### QUESTION NO: 40

A company uses AWS Lambda functions in a private subnet in a VPC to run application logic. The Lambda functions must not have access to the public internet. Additionally, all data communication must remain within the private network. As part of a new requirement, the application logic needs access to an Amazon DynamoDB table.

What is the MOST secure way to meet this new requirement?

- A. Provision the DynamoDB table inside the same VPC that contains the Lambda functions.
- B. Create a gateway VPC endpoint for DynamoDB to provide access to the table.
- C. Use a network ACL to only allow access to the DynamoDB table from the VPC.
- D. Use a security group to only allow access to the DynamoDB table from the VPC.

#### ANSWER: B

##### Explanation:

Create a gateway VPC endpoint for DynamoDB to provide access to the table. A gateway endpoint provides private connectivity from resources in a VPC, including Lambda functions configured with VPC network interfaces, to Amazon DynamoDB. DynamoDB is a regional AWS service and is not deployed inside a customer VPC; instead, the endpoint adds a route to the VPC route tables that directs DynamoDB traffic through the AWS network. This design does not require an internet gateway, NAT gateway, public IP address, or public-internet path, which meets the requirement that the Lambda functions have no public internet access and that application data communications remain private. The endpoint can also be protected with an endpoint policy that limits which DynamoDB actions, tables, and principals may use it. IAM permissions on the Lambda execution role should additionally follow least privilege by allowing only the required operations on the specific table. AWS documents gateway endpoints as the endpoint type for Amazon S3 and Amazon DynamoDB, providing private access without an internet gateway or NAT device. See [AWS VPC endpoint documentation](#) and [Amazon DynamoDB gateway endpoint documentation](#).

#### QUESTION NO: 41

A media streaming company is redesigning its infrastructure to accommodate increasing demand for video content that users consume daily. The company needs to process terabyte-sized videos to block some content in the videos. Video processing can take up to 20 minutes.

The company needs a solution that is cost-effective, highly available, and scalable.

Which solution will meet these requirements?

- A. Use AWS Lambda functions to process the videos. Store video metadata in Amazon DynamoDB. Store video content in Amazon S3 Intelligent-Tiering.
- B. Use Amazon Elastic Container Service (Amazon ECS) with the AWS Fargate launch type to implement microservices to process videos. Store video metadata in Amazon Aurora. Store video content in Amazon S3 Intelligent-Tiering.

C. Use Amazon EMR to process the videos with Apache Spark. Store video content in Amazon FSx for Lustre. Use Amazon Kinesis Data Streams to ingest videos in real time.

D. Deploy a containerized video processing application on Amazon Elastic Kubernetes Service (Amazon EKS) with the Amazon EC2 launch type. Store video metadata in Amazon RDS in a single Availability Zone. Store video content in Amazon S3 Glacier Deep Archive.

**ANSWER: B**

**Explanation:**

Use Amazon Elastic Container Service (Amazon ECS) with the AWS Fargate launch type to implement microservices to process videos. Store video metadata in Amazon Aurora. Store video content in Amazon S3 Intelligent-Tiering. This design supports video-processing jobs that exceed 15 minutes because Fargate tasks can run containerized applications without the AWS Lambda maximum execution-duration constraint. ECS services and tasks can scale horizontally in response to processing demand, while Fargate removes the operational work of provisioning, patching, and managing the underlying compute instances. Video-processing containers can retrieve source objects from Amazon S3, perform the required content-blocking work, and write resulting objects back to S3. Amazon S3 Intelligent-Tiering is appropriate for large media objects when access patterns may change over time, because it automatically moves objects between eligible access tiers to optimize storage cost without operational effort or performance impact for the frequent and infrequent access tiers. Amazon Aurora provides a managed, highly available relational database service for video metadata and can be deployed across multiple Availability Zones. Together, these managed services provide an elastic architecture that can process many video jobs in parallel while maintaining availability and controlling infrastructure-management overhead. See [AWS Fargate for Amazon ECS](#) and [S3 Intelligent-Tiering](#).

**QUESTION NO: 42**

A company hosts multiple applications on AWS for different product lines. The applications use different compute resources, including Amazon EC2 instances and Application Load Balancers. The applications run in different AWS accounts under the same organization in AWS Organizations across multiple AWS Regions. Teams for each product line have tagged each compute resource in the individual accounts.

The company wants more details about the cost for each product line from the consolidated billing feature in Organizations.

Which combination of steps will meet these requirements? (Select TWO.)

- A. Select a specific AWS generated tag in the AWS Billing console.
- B. Select a specific user-defined tag in the AWS Billing console.
- C. Select a specific user-defined tag in the AWS Resource Groups console.
- D. Activate the selected tag from each AWS account.
- E. Activate the selected tag from the Organizations management account.

**ANSWER: B E**

**Explanation:**

**Select a specific user-defined tag in the AWS Billing console.** is correct because the existing product-line tags are user-defined resource tags. To use a tag as a cost dimension in AWS billing tools, AWS must recognize its tag key as a cost allocation tag. After activation and processing, the tag can be selected in AWS Cost Explorer, Cost and Usage Reports, and other Billing and Cost Management views to group, filter, and analyze eligible charges by product line. This supports chargeback or showback reporting across the company's accounts and Regions.

**Activate the selected tag from the Organizations management account.** is correct because centralized cost allocation tag management is performed in the AWS Organizations management account for an organization. Activating the relevant user-defined tag key there makes it available for consolidated cost reporting across member accounts. AWS records activated tags in billing data, although newly activated tags can require up to 24 hours to appear and do not retroactively populate historical cost-allocation data. The organization can then use the common product-line tag key to obtain a consolidated view of costs while retaining the ability to filter by linked account, service, Region, or product line.

### QUESTION NO: 43

A company produces batch data that comes from different databases. The company also produces live stream data from network sensors and application APIs. The company needs to consolidate all the data into one place for business analytics. The company needs to process the incoming data and then stage the data in different Amazon S3 buckets. Teams will later run one-time queries and import the data into a business intelligence tool to show key performance indicators (KPIs).

Which combination of steps will meet these requirements with the LEAST operational overhead? (Choose two.)

- A. Use Amazon Athena for one-time queries. Use Amazon QuickSight to create dashboards for KPIs.
- B. Use Amazon Kinesis Data Analytics for one-time queries Use Amazon QuickSight to create dashboards for KPIs
- C. Create custom AWS Lambda functions to move the individual records from me databases to an Amazon Redshift duster
- D. Use an AWS Glue extract transform, and toad (ETL) job to convert the data into JSON format Load the data into multiple Amazon OpenSearch Service (Amazon Elasticsearch Service) dusters
- E. Use blueprints in AWS Lake Formation to identify the data that can be ingested into a data lake. Use AWS Glue to crawl, transform, and load the data into Amazon S3 in Apache Parquet format.

### ANSWER: A E

#### Explanation:

Using blueprints in AWS Lake Formation to identify data that can be ingested into a data lake, together with AWS Glue to catalog, transform, and load data into Amazon S3 in Apache Parquet format, provides a managed data-lake ingestion and preparation approach. Lake Formation blueprints create repeatable ingestion workflows for data sources, while AWS Glue can perform serverless ETL and maintain metadata in the AWS Glue Data Catalog. Apache Parquet is a columnar format that is well suited to analytical workloads because it can reduce the amount of data scanned and improve query efficiency. Storing the processed results in separate S3 buckets also directly meets the requirement to stage data in Amazon S3. AWS documentation describes Lake Formation blueprints and their workflow-based ingestion capabilities at [AWS Lake Formation blueprints](#).

Using Amazon Athena for one-time queries and Amazon QuickSight for KPI dashboards completes the analytics and visualization workflow without infrastructure to manage. Athena is serverless and queries data directly in Amazon S3 using standard SQL, making it appropriate for ad hoc analysis of the staged Parquet datasets. QuickSight is a managed business intelligence service that can use Athena as a data source and produce interactive analyses and dashboards for KPI reporting. This combination minimizes operational work while allowing teams to query the centralized data lake and present results to business users. See [What is Amazon Athena?](#) for Athena's S3-based serverless query model.

### QUESTION NO: 44

A company hosts an application on AWS. The application gives users the ability to upload photos and store the photos in an Amazon S3 bucket. The company wants to use Amazon CloudFront and a custom domain name to upload the photo files to the S3 bucket in the eu-west-1 Region.

Which solution will meet these requirements? (Select TWO.)

- A. Use AWS Certificate Manager (ACM) to create a public certificate in the us-east-1 Region. Use the certificate in CloudFront.
- B. Use AWS Certificate Manager (ACM) to create a public certificate in eu-west-1. Use the certificate in CloudFront.
- C. Configure Amazon S3 to allow uploads from CloudFront. Configure S3 Transfer Acceleration.
- D. Configure Amazon S3 to allow uploads from CloudFront origin access control (OAC).
- E. Configure Amazon S3 to allow uploads from CloudFront. Configure an Amazon S3 website endpoint.

**ANSWER: A D**

**Explanation:**

Use AWS Certificate Manager (ACM) to create a public certificate in the us-east-1 Region. Use the certificate in CloudFront. This is required because a CloudFront distribution can use an ACM certificate for an alternate domain name only when the certificate is requested or imported in the US East (N. Virginia) Region. The location of the S3 bucket does not change this CloudFront certificate requirement, so the bucket can remain in eu-west-1 while CloudFront serves the custom domain with the certificate from us-east-1. AWS documents this regional requirement for CloudFront viewer certificates at [Requirements for using alternate domain names with CloudFront](#).

Configure Amazon S3 to allow uploads from CloudFront origin access control (OAC). OAC lets CloudFront sign requests to the S3 origin using AWS Signature Version 4. The S3 bucket policy can grant the CloudFront distribution the required object-upload permission, such as `s3:PutObject`, while keeping the bucket private and restricting direct public uploads. The CloudFront cache behavior must also permit the needed upload methods, including PUT or POST as appropriate for the application. AWS describes OAC, signed CloudFront-to-S3 requests, and the required bucket-policy approach at [Restricting access to an Amazon S3 origin](#).

**QUESTION NO: 45**

A company is migrating a business-critical application to AWS. The application must maintain private connectivity between the on-premises data center and a VPC. The company requires consistent high bandwidth and must continue operating if a single AWS Direct Connect location becomes unavailable.

Which combination of actions should a solutions architect take to meet these requirements? (Choose two.)

- A. Provision AWS Direct Connect connections through separate Direct Connect locations.
- B. Configure redundant private virtual interfaces that connect to the VPC through a Direct Connect gateway.
- C. Create an internet gateway in the VPC and use public IP addresses for the application servers.
- D. Create a gateway VPC endpoint for Amazon S3 and route all on-premises traffic through the endpoint.
- E. Use VPC peering to connect the on-premises data center to the VPC.

**ANSWER: A B**

**Explanation:**

Provision AWS Direct Connect connections through separate Direct Connect locations. AWS Direct Connect provides dedicated private network connectivity between an on-premises environment and AWS. Using connections at separate Direct Connect locations removes a single location as a point of failure and supports resilient connectivity for the business-critical workload.

Configure redundant private virtual interfaces (VIFs) to access the VPC through a Direct Connect gateway. A private VIF provides private connectivity to VPC resources through a virtual private gateway or a Direct Connect gateway. The Direct Connect gateway can associate with virtual private gateways in one or more VPCs, and redundant private VIFs support highly available private routing. See the [AWS Direct Connect User Guide](#) and the [AWS Direct Connect virtual interface documentation](#).

**QUESTION NO: 46**

A company is developing a serverless web application that gives users the ability to interact with real-time analytics from online games. The data from the games must be streamed in real time. The company needs a durable, low-latency database option for user data. The company does not know how many users will use the application. Any design considerations must provide response times of single-digit milliseconds as the application scales.

Which combination of AWS services will meet these requirements? Select TWO.

- A. Amazon CloudFront

- B. Amazon DynamoDB
- C. Amazon Kinesis
- D. Amazon RDS
- E. AWS Global Accelerator

**ANSWER: B C**

**Explanation:**

**Amazon Kinesis** provides the managed real-time streaming capability needed to ingest game events continuously and make those records available for real-time processing and analytics. Kinesis Data Streams is designed for collecting and processing streaming data at any scale, and its data records are synchronously replicated across multiple Availability Zones for durability. This makes it an appropriate ingestion layer for high-volume, continuously generated online-game telemetry.

**Amazon DynamoDB** is a serverless, fully managed NoSQL database that delivers single-digit-millisecond performance at virtually any scale. Its on-demand capacity mode is especially suitable when the number of application users and resulting request volume are unknown, because DynamoDB automatically adjusts to changing traffic without capacity planning. DynamoDB also provides durable storage through replication across multiple Availability Zones in an AWS Region. Together, Amazon Kinesis supplies real-time event streaming, while Amazon DynamoDB provides the scalable, low-latency persistence layer for user data required by the application. See the [Amazon Kinesis Data Streams Developer Guide](#) and the [Amazon DynamoDB Developer Guide](#).

**QUESTION NO: 47**

A company uses a three-tier web application to provide training to new employees. The application is accessed for only 12 hours every day. The company is using an Amazon RDS for MySQL DB instance to store information and wants to minimize costs.

What should a solutions architect do to meet these requirements?

- A. Configure an IAM policy for AWS Systems Manager Session Manager. Create an IAM role for the policy. Update the trust relationship of the role. Set up automatic start and stop for the DB instance.
- B. Create an Amazon ElastiCache for Redis cache cluster that gives users the ability to access the data from the cache when the DB instance is stopped. Invalidate the cache after the DB instance is started.
- C. Launch an Amazon EC2 instance. Create an IAM role that grants access to Amazon RDS. Attach the role to the EC2 instance. Configure a cron job to start and stop the EC2 instance on the desired schedule.
- D. Create AWS Lambda functions to start and stop the DB instance. Create Amazon EventBridge (Amazon CloudWatch Events) scheduled rules to invoke the Lambda functions. Configure the Lambda functions as event targets for the rules

**ANSWER: D**

**Explanation:**

Creating AWS Lambda functions to start and stop the DB instance and using Amazon EventBridge scheduled rules to invoke them is the appropriate cost-optimization solution. An Amazon RDS for MySQL DB instance incurs compute charges while it is running. Because the application is predictably unused for approximately 12 hours each day, scheduling the instance to stop outside business hours reduces those compute costs while preserving the database's storage, automated backups, snapshots, parameter configuration, and other persistent instance data.

EventBridge Scheduler or EventBridge scheduled rules can invoke Lambda on a defined daily schedule. The Lambda functions can call the Amazon RDS `StartDBInstance` and `StopDBInstance` API operations, with an IAM execution role granting only the required RDS permissions for the specific DB instance. The start function should run sufficiently before users need the application so that the database is available when the training system opens.

This schedule is compatible with the RDS limitation that a stopped DB instance is automatically started after seven days; the instance is restarted daily, well within that limit. See the [Amazon RDS documentation for stopping and starting DB instances](#) and the [EventBridge documentation for scheduled Lambda invocations](#).

#### QUESTION NO: 48

A solutions architect needs to implement a solution that can handle up to 5,000 messages per second. The solution must publish messages as events to multiple consumers. The messages are up to 500 KB in size. The message consumers need to have the ability to use multiple programming languages to consume the messages with minimal latency. The solution must retain published messages for more than 3 months. The solution must enforce strict ordering of the messages.

- A.** Publish messages to an Amazon Kinesis Data Streams data stream. Enable enhanced fan-out. Ensure that consumers ingest the data stream by using dedicated throughput.
- B.** Publish messages to an Amazon Simple Notification Service (Amazon SNS) topic. Ensure that consumers use an Amazon Simple Queue Service (Amazon SQS) FIFO queue to subscribe to the topic.
- C.** Publish messages to Amazon EventBridge. Allow each consumer to create rules to deliver messages to the consumer's own target.
- D.** Publish messages to an Amazon Simple Notification Service (Amazon SNS) topic. Ensure that consumers use Amazon Data Firehose to subscribe to the topic.

#### ANSWER: A

##### Explanation:

Publish messages to an Amazon Kinesis Data Streams data stream. Enable enhanced fan-out. Ensure that consumers ingest the data stream by using dedicated throughput. This approach meets the event-streaming, scale, retention, fan-out, and low-latency requirements. Kinesis Data Streams accepts records up to 1 MiB, so 500 KB messages are supported. Stream capacity can be scaled by adding shards to accommodate the required ingestion rate. Kinesis retains records for 24 hours by default, and retention can be extended to as long as 365 days, which satisfies retention beyond 3 months.

Enhanced fan-out is particularly appropriate when several consumers need low-latency access to the same stream. Each registered consumer receives dedicated read throughput of up to 2 MiB/second per shard and uses the low-latency HTTP/2 data retrieval path, rather than sharing standard consumer read throughput. Kinesis client libraries and APIs support consumers written in multiple programming languages. Kinesis preserves record order within each shard; therefore, producers should consistently use the same partition key for records that require a strict ordered sequence, ensuring those records are routed to the same shard. See the [Amazon Kinesis Data Streams core concepts](#) and [enhanced fan-out documentation](#).

#### QUESTION NO: 49

A company that uses AWS needs a solution to predict the resources needed for manufacturing processes each month. The solution must use historical values that are currently stored in an Amazon S3 bucket. The company has no machine learning (ML) experience and wants to use a managed service for the training and predictions.

Which combination of steps will meet these requirements? (Select TWO.)

- A.** Deploy an Amazon SageMaker model. Create a SageMaker endpoint for inference.
- B.** Use Amazon SageMaker to train a model by using the historical data in the S3 bucket.
- C.** Configure an AWS Lambda function with a function URL that uses Amazon SageMaker endpoints to create predictions based on the inputs.
- D.** Configure an AWS Lambda function with a function URL that uses an Amazon Forecast predictor to create a prediction based on the inputs.
- E.** Train an Amazon Forecast predictor by using the historical data in the S3 bucket.

**ANSWER: A B****Explanation:**

“Use Amazon SageMaker to train a model by using the historical data in the S3 bucket” and “Deploy an Amazon SageMaker model. Create a SageMaker endpoint for inference.” together provide the required managed training and prediction workflow. Amazon SageMaker can use training data stored in Amazon S3 as the input for a managed training job. SageMaker manages the underlying training infrastructure, including provisioning and operating the compute resources required to train the model. This enables the company to build a monthly resource-prediction model without needing to operate ML infrastructure itself.

After training produces model artifacts, deploying the model to a SageMaker endpoint makes it available through a managed HTTPS inference endpoint. Applications can submit relevant manufacturing inputs to that endpoint and receive prediction results. SageMaker endpoint deployment also provides managed hosting, scalability configuration, monitoring integration, and secure IAM-based access controls. Thus, these steps cover both parts of the requirement: using the historical S3 data to train a model and serving that trained model for ongoing predictions.

For details, see the [Amazon SageMaker training documentation](#) and the [SageMaker real-time inference endpoint documentation](#).

**QUESTION NO: 50**

A telemarketing company is designing its customer call center functionality on AWS. The company needs a solution that provides multiple speaker recognition and generates transcript files. The company wants to query the transcript files to analyze the business patterns. The transcript files must be stored for 7 years for auditing purposes.

Which solution will meet these requirements?

- A. Use Amazon Recognition for multiple speaker recognition. Store the transcript files in Amazon S3. Use machine learning models for transcript file analysis.
- B. Use Amazon Transcribe for multiple speaker recognition. Use Amazon Athena to query transcript file analysis.
- C. Use Amazon Translate for multiple speaker recognition. Store the transcript files in Amazon Redshift. Use SQL queries for transcript file analysis.
- D. Use Amazon Recognition for multiple speaker recognition. Store the transcript files in Amazon S3. Use Amazon Textract for transcript file analysis.
- E. Use Amazon Transcribe with speaker diarization. Store transcript files in Amazon S3 with a 7-year retention period, and use Amazon Athena to query the transcript files.

**ANSWER: E****Explanation:**

Use Amazon Transcribe with speaker diarization, store the transcript files in Amazon S3 with a 7-year retention period, and use Amazon Athena to query the transcript data. Amazon Transcribe is designed to convert speech to text and can identify and label different speakers in an audio stream through speaker diarization. Its transcription output includes speaker labels and timestamps, making it suitable for analyzing call-center conversations involving multiple participants.

Amazon S3 provides durable storage for the generated transcript files. A lifecycle configuration can retain objects for the required period and expire them after seven years. For audit-focused retention requirements where files must not be altered or deleted before the retention date, the company can use S3 Object Lock with an appropriate retention period. Athena can then run SQL queries directly against transcript files in S3, allowing analysts to examine call content, speaker data, timestamps, and other patterns without loading the data into a separate database. See the [Amazon Transcribe speaker diarization documentation](#) and the [Amazon Athena User Guide](#).

**QUESTION NO: 51**

A hospital wants to create digital copies for its large collection of historical written records. The hospital will continue to add hundreds of new documents each day. The hospital's data team will scan the documents and will upload the documents to the AWS Cloud.

A solutions architect must implement a solution to analyze the documents, extract the medical information, and store the documents so that an application can run SQL queries on the data. The solution must maximize scalability and operational efficiency.

Which combination of steps should the solutions architect take to meet these requirements? (Select TWO.)

- A. Write the document information to an Amazon EC2 instance that runs a MySQL database.
- B. Write the document information to an Amazon S3 bucket. Use Amazon Athena to query the data.
- C. Create an Auto Scaling group of Amazon EC2 instances to run a custom application that processes the scanned files and extracts the medical information.
- D. Create an AWS Lambda function that runs when new documents are uploaded. Use Amazon Rekognition to convert the documents to raw text. Use Amazon Transcribe Medical to detect and extract relevant medical information from the text.
- E. Create an AWS Lambda function that runs when new documents are uploaded. Use Amazon Textract to convert the documents to raw text. Use Amazon Comprehend Medical to detect and extract relevant medical information from the text.

**ANSWER: B E**

**Explanation:**

"Create an AWS Lambda function that runs when new documents are uploaded. Use Amazon Textract to convert the documents to raw text. Use Amazon Comprehend Medical to detect and extract relevant medical information from the text." provides a serverless, event-driven processing pipeline. An Amazon S3 upload event can invoke Lambda, which can initiate Amazon Textract processing to extract text from scanned records. The resulting text can then be submitted to Amazon Comprehend Medical, a service designed to identify medical entities and protected health information from unstructured clinical text. This removes the need to operate and scale custom document-processing servers as the daily document volume grows.

"Write the document information to an Amazon S3 bucket. Use Amazon Athena to query the data." provides highly scalable, durable storage for the extracted information and enables SQL-based analysis without provisioning or managing database infrastructure. The extracted results can be stored in an Athena-supported structured format, such as JSON or Parquet, and cataloged for querying. Together, S3, Lambda, Textract, Comprehend Medical, and Athena create an operationally efficient architecture that automatically processes new scans and supports ad hoc SQL queries. See [Amazon Textract documentation](#) and [Amazon Comprehend Medical documentation](#).

**QUESTION NO: 52**

An ecommerce company is planning to migrate an on-premises Microsoft SQL Server database to the AWS Cloud. The company needs to migrate the database to SQL Server Always On availability groups. The cloud-based solution must be highly available.

Options:

- A. Deploy three Amazon EC2 instances with SQL Server across three Availability Zones. Attach one Amazon Elastic Block Store (Amazon EBS) volume to the EC2 instances.
- B. Migrate the database to Amazon RDS for SQL Server. Configure a Multi-AZ deployment and read replicas.
- C. Deploy three Amazon EC2 instances with SQL Server across three Availability Zones. Use Amazon FSx for Windows File Server as the storage tier.
- D. Deploy three Amazon EC2 instances with SQL Server across three Availability Zones. Use Amazon S3 as the storage tier.
- E. Deploy three Amazon EC2 instances with SQL Server across three Availability Zones. Attach separate Amazon EBS volumes to each instance, and configure a Windows Server Failover Cluster and SQL Server Always On availability group.

**ANSWER: E**

**Explanation:**

Deploy three Amazon EC2 instances with SQL Server across three Availability Zones, attach separate Amazon EBS volumes to each instance, and configure a Windows Server Failover Cluster and SQL Server Always On availability group is correct because an availability group consists of separate SQL Server replicas, each maintaining its own copy of the databases and transaction logs. The replicas therefore require independent storage; an Amazon EBS volume is attached to an instance in the same Availability Zone and provides the persistent block storage appropriate for each SQL Server node. Deploying replicas across Availability Zones isolates the workload from an Availability Zone failure. A Windows Server Failover Cluster supplies health detection and quorum, while the availability group provides replication and failover of the SQL Server databases. The workload can use synchronous-commit replicas and automatic failover where the recovery objectives require it. A file share witness, such as one hosted on Amazon FSx for Windows File Server, can be added when needed for quorum design, but it is not shared database storage for an availability group. AWS documents SQL Server high-availability designs on EC2, including Always On availability groups, at [AWS EC2 User Guide](#). Microsoft also describes that each availability-group replica hosts and maintains its own database copy in its [Always On availability groups overview](#).

**QUESTION NO: 53**

A company has multiple VPCs across AWS Regions to support and run workloads that are isolated from workloads in other Regions. Because of a recent application launch requirement, the company's VPCs must communicate with all other VPCs across all Regions.

Which solution will meet these requirements with the LEAST amount of administrative effort?

- A.** Use VPC peering to manage VPC communication in a single Region. Use VPC peering across Regions to manage VPC communications.
- B.** Use AWS Direct Connect gateways across all Regions to connect VPCs across regions and manage VPC communications.
- C.** Use AWS Transit Gateway to manage VPC communication in a single Region and Transit Gateway peering across Regions to manage VPC communications.
- D.** Use AWS PrivateLink across all Regions to connect VPCs across Regions and manage VPC communications.

**ANSWER: C**

**Explanation:**

Use AWS Transit Gateway to manage VPC communication in a single Region and Transit Gateway peering across Regions to manage VPC communications. AWS Transit Gateway provides a centralized network-transit hub to which multiple VPCs can attach. This avoids managing an individual connection between every pair of VPCs, which becomes operationally difficult as the number of VPCs and Regions increases. Route tables on each transit gateway provide centralized control over which attached networks can communicate, while VPC route tables direct relevant traffic to the transit gateway attachment.

For connectivity between Regions, transit gateway peering attachments connect transit gateways in different AWS Regions over the AWS global network. After creating the peering attachment and configuring routes on both transit gateways and their attached VPCs, workloads can communicate privately across the connected regional networks. This hub-and-spoke approach scales far more cleanly than a full mesh of direct VPC-to-VPC connections and supports the requirement for broad VPC connectivity with low administrative overhead. AWS documents Transit Gateway as a service for connecting VPCs and on-premises networks through a central hub, and documents peering attachments specifically for inter-Region transit gateway connectivity. See the [AWS Transit Gateway documentation](#) and [Transit Gateway peering documentation](#).

**QUESTION NO: 54**

A digital image processing company wants to migrate its on-premises monolithic application to the AWS Cloud. The company processes thousands of images and generates large files as part of the processing workflow.

The company needs a solution to manage the growing number of image processing jobs. The solution must also reduce the manual tasks in the image processing workflow. The company does not want to manage the underlying infrastructure of the solution.

Which solution will meet these requirements with the LEAST operational overhead?

- A.** Use Amazon Elastic Container Service (Amazon ECS) with Amazon EC2 Spot Instances to process the images. Configure Amazon Simple Queue Service (Amazon SQS) to orchestrate the workflow. Store the processed files in Amazon Elastic File System (Amazon EFS)
- B.** Use AWS Batch jobs to process the images. Use AWS Step Functions to orchestrate the workflow. Store the processed files in an Amazon S3 bucket.
- C.** Use AWS Lambda functions and Amazon EC2 Spot Instances to process the images. Store the processed files in Amazon FSx.
- D.** Deploy a group of Amazon EC2 instances to process the images. Use AWS Step Functions to orchestrate the workflow. Store the processed files in an Amazon Elastic Block Store (Amazon EBS) volume.

**ANSWER: B**

**Explanation:**

Use AWS Batch jobs to process the images, AWS Step Functions to orchestrate the workflow, and Amazon S3 to store the processed files. AWS Batch is designed for large-scale batch processing and removes the need to provision, configure, and operate the compute fleet that executes jobs. It accepts jobs through queues, schedules them according to their resource requirements and priorities, and automatically provisions the required compute capacity. This directly supports a workload that must scale to thousands of independent image-processing tasks while minimizing infrastructure administration.

AWS Step Functions provides managed workflow orchestration. It can coordinate each stage of image processing, track execution state, apply retries and error handling, and invoke AWS Batch jobs in a controlled sequence. This reduces the manual coordination otherwise required across processing stages. Amazon S3 is a highly scalable, durable object store that is appropriate for the large input and output files generated by the workflow. Together, these managed services separate job execution, workflow coordination, and file storage without requiring the company to manage servers or persistent storage infrastructure. See the [AWS Batch User Guide](#) and the [AWS Step Functions integration with AWS Batch documentation](#).

**QUESTION NO: 55**

A company collects data from thousands of remote devices by using a RESTful web services application that runs on an Amazon EC2 instance. The EC2 instance receives the raw data, transforms the raw data, and stores all the data in an Amazon S3 bucket. The number of remote devices will increase into the millions soon. The company needs a highly scalable solution that minimizes operational overhead.

Which combination of steps should a solutions architect take to meet these requirements? (Select TWO.)

- A.** Use AWS Glue to process the raw data in Amazon S3.
- B.** Use Amazon Route 53 to route traffic to different EC2 instances.
- C.** Add more EC2 instances to accommodate the increasing amount of incoming data.
- D.** Send the raw data to Amazon Simple Queue Service (Amazon SQS). Use EC2 instances to process the data.
- E.** Use Amazon API Gateway to send the raw data to an Amazon Kinesis data stream. Configure Amazon Kinesis Data Firehose to use the data stream as a source to deliver the data to Amazon S3.

**ANSWER: A E**

**Explanation:**

Using Amazon API Gateway to send raw data to an Amazon Kinesis data stream, with Amazon Kinesis Data Firehose delivering that stream to Amazon S3, creates a serverless and highly scalable ingestion path. API Gateway can receive the REST requests from the remote devices without the company operating or scaling a fleet of web-serving EC2 instances.

Kinesis Data Streams is designed for continuously ingesting streaming records at scale, and Firehose can consume records from a Kinesis data stream and reliably deliver them to an S3 bucket. This separates request ingestion from downstream storage and accommodates growth in the device fleet with managed AWS services.

Using AWS Glue to process the raw data in Amazon S3 provides the managed transformation component after ingestion. AWS Glue jobs can read data stored in S3, perform extract, transform, and load processing, and write transformed results back to S3 or other supported data stores. Glue provides managed infrastructure and can scale its processing resources, avoiding the operational work of provisioning, patching, and capacity-planning EC2-based transformation workers. Together, the streaming ingestion and managed ETL design replaces the instance-based workflow while retaining durable S3 storage and supporting substantial increases in incoming device data. See the [Amazon Data Firehose documentation](#) and the [AWS Glue documentation](#).

#### QUESTION NO: 56

A company wants to run its critical applications in containers to meet requirements for scalability and availability. The company prefers to focus on maintenance of the critical applications. The company does not want to be responsible for provisioning and managing the underlying infrastructure that runs the containerized workload.

What should a solutions architect do to meet those requirements?

- A. Use Amazon EC2 Instances, and Install Docker on the Instances
- B. Use Amazon Elastic Container Service (Amazon ECS) on Amazon EC2 worker nodes
- C. Use Amazon Elastic Container Service (Amazon ECS) on AWS Fargate
- D. Use Amazon EC2 instances from an Amazon Elastic Container Service (Amazon ECS)-optimized Amazon Machine Image (AMI).

#### ANSWER: C

##### Explanation:

Use Amazon Elastic Container Service (Amazon ECS) on AWS Fargate. AWS Fargate is a serverless compute engine for containers that lets the company run containerized applications without provisioning, configuring, patching, or managing the Amazon EC2 instances and cluster capacity that host the containers. The company defines the application in an ECS task definition, including the container image, CPU and memory requirements, networking configuration, IAM permissions, and desired task count. ECS then uses Fargate to launch and operate the tasks on AWS-managed infrastructure.

This model aligns with the requirement to concentrate operational effort on the critical applications rather than on the underlying container hosts. For availability and scalability, ECS services can maintain a desired number of healthy tasks, replace failed tasks, and integrate with Elastic Load Balancing for traffic distribution. ECS Service Auto Scaling can adjust the number of running tasks in response to demand or CloudWatch metrics. Fargate also provides per-task resource isolation, allowing workloads to scale without the company having to plan EC2 instance capacity or manage worker-node availability.

AWS describes Fargate as serverless, pay-as-you-go compute for containers that removes the need to manage servers. See the [AWS Fargate for Amazon ECS documentation](#) and the [Amazon ECS Service Auto Scaling documentation](#).

#### QUESTION NO: 57

A company runs a Node.js function on a server in its on-premises data center. The data center stores data in a PostgreSQL database. The company stores the credentials in a connection string in an environment variable on the server. The company wants to migrate its application to AWS and to replace the Node.js application server with AWS Lambda. The company also wants to migrate to Amazon RDS for PostgreSQL and to ensure that the database credentials are securely managed.

Which solution will meet these requirements with the LEAST operational overhead?

- A. Store the database credentials as a parameter in AWS Systems Manager Parameter Store. Configure Parameter Store to automatically rotate the secrets every 30 days. Update the Lambda function to retrieve the credentials from the parameter.

- B.** Store the database credentials as a secret in AWS Secrets Manager. Configure Secrets Manager to automatically rotate the credentials every 30 days. Update the Lambda function to retrieve the credentials from the secret.
- C.** Store the database credentials as an encrypted Lambda environment variable. Write a custom Lambda function to rotate the credentials. Schedule the Lambda function to run every 30 days.
- D.** Store the database credentials as a key in AWS Key Management Service (AWS KMS). Configure automatic rotation for the key. Update the Lambda function to retrieve the credentials from the KMS key.

**ANSWER: B**

**Explanation:**

Store the database credentials as a secret in AWS Secrets Manager. Configure Secrets Manager to automatically rotate the credentials every 30 days. Update the Lambda function to retrieve the credentials from the secret. This solution uses the AWS service purpose-built for managing sensitive application credentials, including credentials for Amazon RDS for PostgreSQL. Secrets Manager encrypts the secret at rest using AWS Key Management Service and lets the Lambda function retrieve the current credential securely at runtime through the AWS SDK, controlled by an IAM execution-role policy.

Secrets Manager also provides managed rotation workflows for supported Amazon RDS databases. When rotation is configured, Secrets Manager invokes a rotation Lambda function according to the defined schedule, updates the database credential, validates it, and maintains the current secret value for consumers. This avoids embedding a static connection string in application configuration and minimizes the custom code and operational processes needed to change credentials regularly. The Lambda function should retrieve and, where appropriate, cache the secret value, then establish its PostgreSQL connection using the current username and password. See the [AWS Secrets Manager rotation documentation](#) and [retrieving secrets in AWS Lambda documentation](#).

**QUESTION NO: 58**

A company is developing a new online gaming application. The application will run on Amazon EC2 instances in multiple AWS Regions and will have a high number of globally distributed users. A solutions architect must design the application to optimize network latency for the users.

Which actions should the solutions architect take to meet these requirements? (Select TWO.)

- A.** Configure AWS Global Accelerator. Create Regional endpoint groups in each Region where an EC2 fleet is hosted.
- B.** Create a content delivery network (CDN) by using Amazon CloudFront. Enable caching for static and dynamic content, and specify a high expiration period.
- C.** Integrate AWS Client VPN into the application. Instruct users to select which Region is closest to them after they launch the application. Establish a VPN connection to that Region.
- D.** Create an Amazon Route 53 weighted routing policy. Configure the routing policy to give the highest weight to the EC2 instances in the Region that has the largest number of users.
- E.** Configure an Amazon API Gateway endpoint in each Region where an EC2 fleet is hosted. Instruct users to select which Region is closest to them after they launch the application. Use the API Gateway endpoint that is closest to them.

**ANSWER: A B**

**Explanation:**

**Configure AWS Global Accelerator. Create Regional endpoint groups in each Region where an EC2 fleet is hosted.** is correct because Global Accelerator provides globally anycast static IP addresses and routes player traffic onto the AWS global network at the edge location nearest to the user. It continuously evaluates endpoint health and directs connections to the optimal healthy Regional endpoint, reducing the latency and variability associated with traversing the public internet. Regional endpoint groups allow the EC2-hosted gaming fleet in each AWS Region to participate in this global traffic-routing design. AWS documents this behavior at [What is AWS Global Accelerator?](#).

**Create a content delivery network (CDN) by using Amazon CloudFront. Enable caching for static and dynamic content, and specify a high expiration period.** is also correct because CloudFront uses a worldwide network of edge

locations to deliver cacheable game assets, such as downloads, images, scripts, patches, and web content, from locations close to players. CloudFront can also improve delivery for dynamic content by using optimized persistent connections and the AWS network between edge locations and origins. Caching appropriate content at the edge reduces requests to the Regional EC2 fleets and improves response time for globally distributed users. CloudFront caching behavior and TTL controls are described in the [Amazon CloudFront Developer Guide](#).

#### QUESTION NO: 59

A company stores critical data in Amazon DynamoDB tables in the company's AWS account. An IT administrator accidentally deleted a DynamoDB table. The deletion caused a significant loss of data and disrupted the company's operations. The company wants to prevent this type of disruption in the future.

Which solution will meet this requirement with the LEAST operational overhead?

- A. Configure a trail in AWS CloudTrail. Create an Amazon EventBridge rule for delete actions. Create an AWS Lambda function to automatically restore deleted DynamoDB tables.
- B. Create a backup and restore plan for the DynamoDB tables. Recover the DynamoDB tables manually.
- C. Configure deletion protection on the DynamoDB tables.
- D. Enable point-in-time recovery on the DynamoDB tables.

#### ANSWER: C

##### Explanation:

Configure deletion protection on the DynamoDB tables. DynamoDB deletion protection is specifically designed to prevent accidental table deletion. When this setting is enabled, a request to delete the table is rejected unless an authorized user first explicitly disables deletion protection. This creates an intentional, auditable extra step before a destructive operation can occur, directly preventing the operational disruption described in the scenario.

Deletion protection is a table-level setting that can be enabled through the AWS Management Console, AWS CLI, SDKs, or infrastructure as code. It does not require the company to operate recovery workflows, event processing, custom automation, or periodic backup processes. As a result, it provides the requested preventive control with minimal ongoing administration. The setting applies whether a deletion is attempted through the console, API, CLI, or an AWS service acting on behalf of an appropriately authorized principal. AWS documents that deletion protection must be disabled before a protected table can be deleted, making it an effective guardrail for critical DynamoDB tables. For implementation details, see the [DynamoDB deletion protection documentation](#) and the [UpdateTable API reference](#).

#### QUESTION NO: 60

A company runs a critical public application on Amazon Elastic Kubernetes Service (Amazon EKS) clusters. The application has a microservices architecture. The company needs to implement a solution that collects, aggregates, and summarizes metrics and logs from the application in a centralized location.

Which solution will meet these requirements in the MOST operationally efficient way?

- A. Run the Amazon CloudWatch agent in the existing EKS cluster. Use a CloudWatch dashboard to view the metrics and logs.
- B. Configure a data stream in Amazon Kinesis Data Streams. Use Amazon Kinesis Data Firehose to read events and to deliver the events to an Amazon S3 bucket. Use Amazon Athena to view the events.
- C. Configure AWS CloudTrail to capture data events. Use Amazon OpenSearch Service to query CloudTrail.
- D. Configure Amazon CloudWatch Container Insights in the existing EKS cluster. Use a CloudWatch dashboard to view the metrics and logs.

#### ANSWER: D

### Explanation:

Configure Amazon CloudWatch Container Insights in the existing EKS cluster. Use a CloudWatch dashboard to view the metrics and logs. Container Insights is purpose-built for observability of containerized workloads, including Amazon EKS. It collects, aggregates, and summarizes operational metrics at the cluster, node, pod, container, and service levels, enabling teams to monitor the health and performance of a microservices application from a centralized CloudWatch location. For logs, Container Insights integrates with CloudWatch Logs to collect container log output, where it can be searched, retained, and analyzed alongside metrics. CloudWatch dashboards then provide a managed visualization layer for reviewing this telemetry without operating a separate logging or metrics platform. This is operationally efficient because AWS provides the integration and managed backend; the implementation primarily involves deploying the supported observability components and configuring the EKS cluster. It also supports CloudWatch alarms and Logs Insights queries, allowing the company to build monitoring and investigation workflows around the same centralized data. AWS documents Container Insights specifically as a solution for collecting metrics and logs from containerized applications and infrastructure, including EKS. See the [Amazon CloudWatch Container Insights documentation](#) and the [Amazon EKS Container Insights setup guide](#).

### QUESTION NO: 61

A hospital needs to store patient records in an Amazon S3 bucket. The hospital's compliance team must ensure that all protected health information (PHI) is encrypted in transit and at rest. The compliance team must administer the encryption key for data at rest.

Which solution will meet these requirements?

- A.** Create a public SSL/TLS certificate in AWS Certificate Manager (ACM). Associate the certificate with Amazon S3. Configure default encryption for each S3 bucket to use server-side encryption with AWS KMS keys (SSE-KMS). Assign the compliance team to manage the KMS keys.
- B.** Use the `aws:SecureTransport` condition on S3 bucket policies to allow only encrypted connections over HTTPS (TLS). Configure default encryption for each S3 bucket to use server-side encryption with S3 managed encryption keys (SSE-S3). Assign the compliance team to manage the SSE-S3 keys.
- C.** Use the `aws:SecureTransport` condition on S3 bucket policies to allow only encrypted connections over HTTPS (TLS). Configure default encryption for each S3 bucket to use server-side encryption with AWS KMS keys (SSE-KMS). Assign the compliance team to manage the KMS keys.
- D.** Use the `aws:SecureTransport` condition on S3 bucket policies to allow only encrypted connections over HTTPS (TLS). Use Amazon Macie to protect the sensitive data that is stored in Amazon S3. Assign the compliance team to manage Macie.

### ANSWER: C

### Explanation:

Use the `aws:SecureTransport` condition on S3 bucket policies to allow only encrypted connections over HTTPS (TLS). Configure default encryption for each S3 bucket to use server-side encryption with AWS KMS keys (SSE-KMS). Assign the compliance team to manage the KMS keys. This solution addresses both required encryption states and provides the compliance team with administrative control of the key used to protect data at rest. An S3 bucket policy can explicitly deny requests when the `aws:SecureTransport` condition is false, ensuring that clients use HTTPS and TLS rather than unencrypted HTTP for S3 access. For stored objects, default bucket encryption with SSE-KMS causes Amazon S3 to encrypt new objects with a customer managed AWS KMS key. The compliance team can then administer that KMS key through its key policy and IAM permissions, including controlling who may use it, manage its policy, rotate key material where applicable, and review its use through AWS logging services. SSE-KMS integrates this customer-controlled key management with S3's server-side encryption workflow, so applications do not need to implement their own encryption process before upload. AWS documents the `aws:SecureTransport` global condition key for enforcing TLS and describes SSE-KMS as encryption using AWS KMS keys. See [Amazon S3 policy condition keys](#) and [Using SSE-KMS with Amazon S3](#).

### QUESTION NO: 62

A company has an AWS Lambda function and an Amazon S3 bucket. A solutions architect creates an IAM role that has `S3:GetObject` and `S3:ListBucket` permissions and configures it as the Lambda function execution role. The function must write logs to an Amazon CloudWatch Logs log group when the function is invoked.

Which solution will meet this requirement?

- A. Create a new IAM role for the S3 bucket and grant the role the logs:CreateLogGroup, logs:CreateLogStream, and logs:PutLogEvents permissions.
- B. Update the IAM policy that is attached to the existing IAM role to include the logs:CreateLogGroup, logs:CreateLogStream, and logs:PutLogEvents permissions.
- C. Create an S3 bucket policy that allows the existing IAM role to access the bucket and to write logs to the bucket.
- D. Update the existing IAM role to attach the logs:GetLogEvents IAM policy.

**ANSWER: B**

**Explanation:**

Update the IAM policy that is attached to the existing IAM role to include the `logs:CreateLogGroup`, `logs:CreateLogStream`, and `logs:PutLogEvents` permissions. Lambda assumes its configured execution role every time it runs, so the permissions needed for the function's AWS API calls—including its delivery of invocation logs to CloudWatch Logs—must be available through that role. CloudWatch Logs uses `logs:CreateLogGroup` when a log group must be created, `logs:CreateLogStream` to create a stream for the function environment, and `logs:PutLogEvents` to publish log events. AWS provides these permissions in the managed `AWSLambdaBasicExecutionRole` policy; they can also be added as equivalent permissions in a customer-managed policy attached to the function's existing execution role. Retaining the S3 permissions on that execution role allows the function to continue accessing its required S3 resources while adding only the logging permissions necessary to meet the requirement. For AWS guidance on Lambda execution-role permissions and CloudWatch Logs access, see the [AWS Lambda execution role documentation](#) and the [AWSLambdaBasicExecutionRole policy reference](#).

**QUESTION NO: 63**

A company operates a serverless order-processing application. Several downstream systems need to receive order events, but each system needs only events that match specific attributes. The company needs a solution that can route events by content and retain events that cannot be delivered successfully to a target.

Which combination of actions should a solutions architect take to meet these requirements? (Choose two.)

- A. Create an Amazon EventBridge custom event bus and publish the order events to the event bus.
- B. Create EventBridge rules that use event patterns to route matching events to each downstream target. Configure a dead-letter queue for each target.
- C. Create one Amazon SQS queue and configure every downstream system to consume messages from the queue.
- D. Store each order event in an Amazon S3 bucket. Configure an S3 Lifecycle rule to route events to downstream systems.
- E. Create an Amazon SNS topic and require each downstream system to filter messages by using application code.

**ANSWER: A B**

**Explanation:**

Create an Amazon EventBridge custom event bus and publish order events to the event bus. Amazon EventBridge is designed for event-driven architectures and can receive custom application events. Event producers remain decoupled from downstream consumers because producers publish events without needing to know which services will process them.

Create EventBridge rules with event patterns that filter events based on order attributes, and configure a dead-letter queue for each rule target. Event patterns allow EventBridge to send only matching events to each downstream target. A dead-letter queue captures events that cannot be delivered successfully after EventBridge retry attempts, allowing the company to investigate and replay failed events. See the [Amazon EventBridge event bus documentation](#) and the [EventBridge dead-letter queue documentation](#).

#### QUESTION NO: 64

A solutions architect is designing the architecture for a software demonstration environment. The environment will run on Amazon EC2 instances in an Auto Scaling group behind an Application Load Balancer (ALB). The system will experience significant increases in traffic during working hours but is not required to operate on weekends.

Which combination of actions should the solutions architect take to ensure that the system can scale to meet demand? (Select TWO)

- A. Use AWS Auto Scaling to adjust the ALB capacity based on request rate
- B. Use AWS Auto Scaling to scale the capacity of the VPC internet gateway
- C. Launch the EC2 instances in multiple AWS Regions to distribute the load across Regions
- D. Use a target tracking scaling policy to scale the Auto Scaling group based on instance CPU utilization
- E. Use scheduled scaling to change the Auto Scaling group minimum, maximum, and desired capacity to zero for weekends. Revert to the default values at the start of the week

**ANSWER: D E**

#### Explanation:

Use a target tracking scaling policy to scale the Auto Scaling group based on instance CPU utilization because this policy automatically adjusts the number of EC2 instances to maintain a specified average CPU-utilization target across the group. As traffic grows during working hours, CPU utilization rises and the Auto Scaling group launches additional healthy instances; as demand falls, it can remove excess capacity. This provides responsive, metric-driven scaling while the Application Load Balancer distributes requests across the registered instances.

Use scheduled scaling to change the Auto Scaling group minimum, maximum, and desired capacity to zero for weekends. Revert to the default values at the start of the week because the workload has a known, recurring period when it does not need to operate. Scheduled scaling can set the group's desired, minimum, and maximum capacity to zero for the weekend, avoiding EC2 compute costs. A scheduled action before business hours can restore the normal capacity boundaries and desired capacity, allowing the group to serve expected weekday traffic. Target tracking then continues to adapt capacity to actual load during those operating periods. AWS documents target tracking and scheduled scaling for EC2 Auto Scaling at [Target tracking scaling policies](#) and [Scheduled scaling](#).

#### QUESTION NO: 65

A company has multiple AWS accounts that use consolidated billing. The company runs several active high performance Amazon RDS for Oracle On-Demand DB instances

for 90 days. The company's finance team has access to AWS Trusted Advisor in the consolidated billing account and all other AWS accounts.

The finance team needs to use the appropriate AWS account to access the Trusted Advisor check recommendations for RDS. The finance team must review the

appropriate Trusted Advisor check to reduce RDS costs.

Which combination of steps should the finance team take to meet these requirements? (Select TWO.)

- A. Use the Trusted Advisor recommendations from the account where the RDS instances are running.
- B. Use the Trusted Advisor recommendations from the consolidated billing account to see all RDS instance checks at the same time.
- C. Review the Trusted Advisor check for Amazon RDS Reserved Instance Optimization.
- D. Review the Trusted Advisor check for Amazon RDS Idle DB Instances.
- E. Review the Trusted Advisor check for Amazon Redshift Reserved Node Optimization.

**ANSWER: A C****Explanation:**

Use the Trusted Advisor recommendations from the account where the RDS instances are running because Trusted Advisor recommendations and check results are generated for the individual AWS account that owns the resources. Consolidated billing centralizes billing and can provide aggregate cost-management capabilities, but it does not inherently make an account's Trusted Advisor resource-level check results available in the consolidated billing account. Since the finance team has Trusted Advisor access in every account, it can sign in to, or assume an appropriate role in, each workload account to review the RDS recommendations for the DB instances in that account. AWS documents Trusted Advisor as providing checks and recommendations for an AWS account's deployed resources: [AWS Trusted Advisor User Guide](#).

Review the Trusted Advisor check for Amazon RDS Reserved Instance Optimization because the company has active, stable Amazon RDS for Oracle On-Demand DB instances that have run for 90 days. This check identifies RDS usage for which Reserved DB Instances could provide lower costs than continued On-Demand pricing. Reserved DB Instances are a suitable cost-optimization mechanism for predictable, long-running database workloads, while retaining the operational model of Amazon RDS. The check therefore directly helps the finance team identify reservation purchase opportunities for the existing Oracle database fleet. AWS lists Reserved Instance optimization among its cost-optimizing Trusted Advisor checks: [Cost optimization checks](#).

**QUESTION NO: 66**

A company has an application that ingests incoming messages. These messages are then quickly consumed by dozens of other applications and microservices.

The number of messages varies drastically and sometimes spikes as high as 100,000 each second. The company wants to decouple the solution and increase scalability.

Which solution meets these requirements?

- A.** Persist the messages to Amazon Kinesis Data Analytics. All the applications will read and process the messages.
- B.** Deploy the application on Amazon EC2 instances in an Auto Scaling group, which scales the number of EC2 instances based on CPU metrics.
- C.** Write the messages to Amazon Kinesis Data Streams with a single shard. All applications will read from the stream and process the messages.
- D.** Publish the messages to an Amazon Simple Notification Service (Amazon SNS) topic with one or more Amazon Simple Queue Service (Amazon SQS) subscriptions. All applications then process the messages from the queues.

**ANSWER: D****Explanation:**

Publishing the messages to an Amazon Simple Notification Service (Amazon SNS) topic with one or more Amazon Simple Queue Service (Amazon SQS) subscriptions is the appropriate architecture because it provides a managed fan-out pattern. The producer publishes each incoming message only once to the SNS topic, and SNS independently delivers a copy to every subscribed SQS queue. Each application or microservice can therefore consume messages from its own queue at its own rate, without being coupled to the producer or to other consumers.

SQS acts as a durable buffer during spikes, allowing consumers to scale independently and process accumulated work asynchronously. Standard SQS queues support very high throughput and are designed to scale elastically, which suits unpredictable traffic such as bursts of 100,000 messages per second. Consumers can use metrics such as queue depth and age of the oldest message to drive AWS Auto Scaling decisions. Separate queues also provide workload isolation: a temporary slowdown in one consuming service does not prevent the other services from receiving and processing their copies of the messages.

AWS documents this SNS-to-SQS design as a fan-out messaging pattern for distributing messages to multiple independent consumers. See the [Amazon SNS common messaging scenarios documentation](#) and the [Amazon SQS standard queues documentation](#).

## QUESTION NO: 67

A company runs an order management application on AWS. The application allows customers to place orders and pay with a credit card. The company uses an Amazon CloudFront distribution to deliver the application. A security team has set up logging for all incoming requests. The security team needs a solution to generate an alert if any user modifies the logging configuration.

Which combination of solutions will meet these requirements? (Select TWO.)

- A.** Configure an Amazon EventBridge rule that is invoked when a user creates or modifies a CloudFront distribution. Add the AWS Lambda function as a target of the EventBridge rule.
- B.** Create an Application Load Balancer (ALB). Enable AWS WAF rules for the ALB. Configure an AWS Config rule to detect security violations.
- C.** Create an AWS Lambda function to detect changes in CloudFront distribution logging. Configure the Lambda function to use Amazon Simple Notification Service (Amazon SNS) to send notifications to the security team.
- D.** Set up Amazon GuardDuty. Configure GuardDuty to monitor findings from the CloudFront distribution. Create an AWS Lambda function to address the findings.
- E.** Create a private API in Amazon API Gateway. Use AWS WAF rules to protect the private API from common security problems.

**ANSWER: A C**

### Explanation:

Configuring an Amazon EventBridge rule that is invoked when a user creates or modifies a CloudFront distribution, with an AWS Lambda function as the target, provides the event-driven detection mechanism. CloudFront distribution configuration changes are performed through CloudFront API operations, such as `UpdateDistribution`. AWS CloudTrail records these management events, and Amazon EventBridge can match CloudTrail API events based on details such as the service, API operation, and caller. The rule can be scoped specifically to distribution-update events so the function runs promptly when a configuration change occurs.

Creating an AWS Lambda function to detect changes in CloudFront distribution logging and configuring it to publish notifications through Amazon SNS completes the alerting workflow. The function can inspect the CloudTrail event and, when needed, retrieve the updated distribution configuration to evaluate the logging settings. It can then publish a clear alert to an SNS topic subscribed to by the security team, such as through email, SMS, or an incident-management integration. This combination provides both detection of the configuration change and notification to the appropriate responders. See [Amazon EventBridge integration with AWS CloudTrail events](#) and the [CloudFront UpdateDistribution API reference](#).

## QUESTION NO: 68

A company runs multiple workloads in separate AWS environments. The company wants to optimize its AWS costs but must maintain the same level of performance for the environments.

The company 's production environment requires resources to be highly available. The other environments do not require highly available resources.

Each environment has the same set of networking components, including the following:

- 1 VPC
- 1 Application Load Balancer
- 4 subnets distributed across 2 Availability Zones 2 public subnets and 2 private subnets
- 2 NAT gateways 1 in each public subnet
- 1 internet gateway

Which solution will meet these requirements?

- A.** Do not change the production environment workload. For each non-production workload, remove one NAT gateway and update the route tables for private subnets to target the remaining NAT gateway for the destination 0.0.0.0/0.
- B.** Reduce the number of Availability Zones that all workloads in all environments use.
- C.** Replace every NAT gateway with a t4g.large NAT instance. Update the route tables for each private subnet to target the NAT instance that is in the same Availability Zone for the destination 0.0.0.0/0.
- D.** In each environment, create one transit gateway and remove one NAT gateway.

**ANSWER: A**

**Explanation:**

Do not change the production environment workload. For each non-production workload, remove one NAT gateway and update the route tables for private subnets to target the remaining NAT gateway for the destination 0.0.0.0/0. This solution aligns the architecture and its cost with each environment's availability requirement. A NAT gateway enables instances in private subnets to initiate outbound connections to the internet or AWS services while preventing unsolicited inbound internet connections. In a highly available production VPC spanning multiple Availability Zones, AWS recommends deploying a NAT gateway in each Availability Zone and routing each private subnet to the NAT gateway in its own Availability Zone. This avoids a single point of failure for outbound connectivity and avoids relying on cross-AZ routing during an Availability Zone disruption.

Because the non-production environments do not require highly available resources, they can use one remaining NAT gateway while both private-subnet route tables send their default route, 0.0.0.0/0, to it. The private workloads retain the same outbound internet access functionality and performance characteristics needed for normal operation, while the company eliminates the hourly charge and associated processing charges for the removed NAT gateway. This targets the stated cost optimization goal without reducing the production environment's resilience. AWS documents NAT gateway availability guidance and routing behavior at [NAT gateways](#) and [NAT gateway scenarios](#).

**QUESTION NO: 69**

A company runs a mobile game app on AWS. The app stores data for every user session. The data updates frequently during a gaming session. The app stores up to 256 KB for each session. Sessions can last up to 48 hours.

The company wants to automate the deletion of expired session data. The company must be able to restore all session data automatically if necessary.

Which solution will meet these requirements?

- A.** Use an Amazon DynamoDB table to store the session data. Enable point-in-time recovery (PITR) and TTL for the table. Select the corresponding attribute for TTL in the session data.
- B.** Use an Amazon MemoryDB table to store the session data. Enable point-in-time recovery (PITR) and TTL for the table. Select the corresponding attribute for TTL in the session data.
- C.** Store session data in an Amazon S3 bucket. Use the S3 Standard storage class. Enable S3 Versioning for the bucket. Create an S3 Lifecycle configuration to expire objects after 48 hours.
- D.** Store session data in an Amazon S3 bucket. Use the S3 Intelligent-Tiering storage class. Enable S3 Versioning for the bucket. Create an S3 Lifecycle configuration to expire objects after 48 hours.

**ANSWER: A**

**Explanation:**

Use an Amazon DynamoDB table to store the session data. Enable point-in-time recovery (PITR) and TTL for the table. Select the corresponding attribute for TTL in the session data. DynamoDB is well suited to high-frequency, low-latency updates of session-state records, and its maximum item size is 400 KB, which accommodates the stated session-data size of up to 256 KB. The application can write an expiration timestamp to a designated TTL attribute for each session. DynamoDB TTL then automatically removes expired items without requiring the company to run deletion jobs or manage

infrastructure. TTL deletion is asynchronous, so applications should treat the timestamp as the authoritative expiration point when deciding whether a session remains valid.

Point-in-time recovery provides continuous backups and enables restoration of the entire DynamoDB table to any second within the previous 35 days. This satisfies the requirement to recover session data if deletion or another operational issue requires a restore. PITR is enabled at the table level and does not affect the application's normal read and write behavior. Together, TTL handles automated expiration while PITR provides the required recoverability for the table's data. See the [DynamoDB TTL documentation](#) and the [DynamoDB point-in-time recovery documentation](#).

#### QUESTION NO: 70

A company is performing a security review of its Amazon EMR API usage. The company's developers use an integrated development environment (IDE) that is hosted on Amazon EC2 instances. The IDE is configured to authenticate users to AWS by using access keys. Traffic between the company's EC2 instances and EMR cluster uses public IP addresses.

A solutions architect needs to improve the company's overall security posture. The solutions architect needs to reduce the company's use of long-term credentials and to limit the amount of communication that uses public IP addresses.

Which combination of steps will MOST improve the security of the company's architecture? (Select TWO.)

- A. Set up a gateway endpoint to the EMR cluster.
- B. Set up interface VPC endpoints to connect to the Amazon EMR API.
- C. Set up a private NAT gateway to connect to the EMR cluster.
- D. Set up IAM roles for the developers to use to connect to the Amazon EMR API.
- E. Set up AWS Systems Manager Parameter Store to store access keys for each developer.

#### ANSWER: B D

##### Explanation:

Set up interface VPC endpoints to connect to the Amazon EMR API and set up IAM roles for the developers to use to connect to the Amazon EMR API. An interface VPC endpoint uses AWS PrivateLink to create private network interfaces in the VPC for supported AWS services, including Amazon EMR API endpoints. Requests from the EC2-hosted IDE can therefore reach the EMR service API through private IP addresses within the AWS network rather than traversing public IP addressing. Endpoint security groups and endpoint policies can further constrain which sources and API actions are allowed.

IAM roles address the long-term-credential requirement because roles provide temporary security credentials through AWS Security Token Service. The EC2 instances hosting the IDE can receive permissions through an instance profile, allowing AWS SDKs and the AWS CLI to retrieve and rotate credentials automatically instead of storing developers' long-lived access keys. The role's IAM policy should grant only the EMR actions and resources required by the developers, applying least privilege. Together, private API connectivity and temporary role credentials reduce both public network exposure and the operational risk associated with static access keys. See [Amazon EMR interface VPC endpoints](#) and [IAM roles for Amazon EC2](#).

#### QUESTION NO: 71

A company hosts an internal service on Amazon EC2 instances behind a Network Load Balancer in a provider VPC. Several application teams in separate AWS accounts need private access to the service. The company does not want to create VPC peering connections or expose the service to the public internet.

Which combination of actions should a solutions architect take to meet these requirements? (Choose two.)

- A. Create a VPC endpoint service that is associated with the Network Load Balancer.
- B. Create interface VPC endpoints for the endpoint service in the consumer VPCs.
- C. Create a gateway VPC endpoint in each consumer VPC.

D. Create VPC peering connections between the provider VPC and each consumer VPC.

E. Assign public IP addresses to the Network Load Balancer targets and allow inbound traffic from the consumer VPC CIDR ranges.

**ANSWER: A B**

**Explanation:**

Creating a VPC endpoint service and associating it with the Network Load Balancer makes the provider service available through AWS PrivateLink. The endpoint service can be configured to allow connections only from approved AWS principals or accounts.

Creating interface VPC endpoints in the consumer VPCs provides private network interfaces through which the application teams can access the endpoint service. Traffic remains on the AWS network, and the provider and consumer VPCs do not require VPC peering, transit gateway connectivity, or internet routing. See the [AWS PrivateLink endpoint service documentation](#) and the [interface VPC endpoint documentation](#).

**QUESTION NO: 72**

A company runs a web application that uses Amazon RDS for MySQL to store relational data. Data in the database does not change frequently.

A solutions architect notices that during peak usage times, the database has performance issues when it serves the data. The company wants to improve the performance of the database.

Which combination of steps will meet these requirements? (Select TWO.)

A. Integrate AWS WAF with the application.

B. Create a read replica for the database. Redirect read traffic to the read replica.

C. Create an Amazon ElastiCache (Memcached) cluster. Configure the application to use the cluster as a cache.

D. Use the Amazon S3 One Zone-Infrequent Access (S3 One Zone-IA) storage class to store the data that changes infrequently.

E. Migrate the database to Amazon DynamoDB. Configure the application to use the DynamoDB database.

**ANSWER: B C**

**Explanation:**

**Create a read replica for the database. Redirect read traffic to the read replica.** is appropriate because Amazon RDS for MySQL read replicas asynchronously replicate data from the source DB instance and can serve read-only workloads. Routing read queries from the application to one or more replicas offloads the primary DB instance, allowing it to retain capacity for writes and reducing contention during peak demand. This is especially suitable for an application whose workload includes substantial reads of relational data.

**Create an Amazon ElastiCache (Memcached) cluster. Configure the application to use the cluster as a cache.** further improves performance by keeping frequently requested data in low-latency, in-memory storage. Because the database data changes infrequently, cached results can remain useful for longer periods with an appropriate expiration policy. The application can use a cache-aside pattern: check Memcached first, retrieve a missing item from RDS, then populate the cache. Combining a read replica with ElastiCache reduces repeated database queries while scaling reads independently from the writer database.

For implementation guidance, see the [Amazon RDS read replica documentation](#) and the [ElastiCache caching strategies documentation](#).

**QUESTION NO: 73**

A company is developing a social media application that must scale rapidly and handle long-running, ordered processes that store large amounts of relational data. Components must scale independently and evolve without downtime.

Which combination of AWS services will meet these requirements?

- A. Amazon ECS with Fargate, Amazon RDS, and Amazon SQS FIFO
- B. Amazon ECS with Fargate, Amazon RDS, and Amazon SNS
- C. AWS Lambda, Amazon DynamoDB Streams, and AWS Step Functions
- D. AWS Elastic Beanstalk, Amazon RDS, and Amazon SNS

**ANSWER: A**

**Explanation:**

**Amazon ECS with Fargate, Amazon RDS, and Amazon SQS FIFO** meets the application's requirements by separating the application into independently deployable and scalable containerized services. Amazon ECS on AWS Fargate runs containers without requiring the company to provision or manage EC2 instances. Each ECS service can have its own deployment process and Service Auto Scaling policy, allowing individual components to scale according to their own demand and be updated with rolling deployments while maintaining availability.

Amazon RDS is a managed relational database service suited to storing structured, relational application data. It supports operational capabilities such as backups, Multi-AZ deployments, read replicas for supported engines, and scaling choices that help support a growing application. For the asynchronous long-running work, an Amazon SQS FIFO queue decouples producers from workers and preserves message order within each message group. FIFO queues are specifically designed for ordered processing and provide exactly-once processing behavior, making them appropriate when process sequence matters. ECS worker services can consume queued jobs at a rate that can be scaled independently from the user-facing services.

See the [Amazon SQS FIFO queues documentation](#) and the [Amazon ECS service auto scaling documentation](#).

#### QUESTION NO: 74

A company has a video editing application that requires consistent sub-millisecond latency and high throughput to access media objects that are updated frequently. The company currently has an Amazon S3 bucket that uses the S3 Standard storage class.

The company needs to improve performance while maintaining Amazon S3 API compatibility. The company needs to access the media objects within a single Availability Zone.

Which storage solution will meet these requirements?

- A. Enable S3 Transfer Acceleration. Configure the application to use multipart uploads.
- B. Create an S3 directory bucket that uses the S3 Express One Zone storage class.
- C. Create an S3 table bucket and table namespace.
- D. Use the S3 One Zone-Infrequent Access (S3 One Zone-IA) storage class. Configure the application to use multipart uploads.

**ANSWER: B**

**Explanation:**

**Create an S3 directory bucket that uses the S3 Express One Zone storage class.** S3 Express One Zone is purpose-built for performance-sensitive workloads that need consistent single-digit millisecond latency, and it can provide latency as low as single-digit milliseconds for object access, along with very high request rates. It is particularly suitable for interactive media processing and video editing workflows in which objects are accessed and modified frequently. Unlike storage choices designed primarily for archival or infrequently accessed data, S3 Express One Zone is optimized for active data and high-throughput operations.

S3 Express One Zone stores data in a single Availability Zone, directly satisfying the stated single-AZ requirement. It is used through an S3 directory bucket, which supports an S3-compatible object storage model and enables applications to continue using Amazon S3 APIs while benefiting from the performance characteristics of the Express One Zone storage class. Directory buckets also organize objects hierarchically, helping applications efficiently work with large media asset collections. For implementation details and supported API behavior, see the [Amazon S3 directory buckets documentation](#) and the [S3 Express One Zone documentation](#).

#### QUESTION NO: 75

A company operates a data lake in Amazon S3. The company wants to query and filter data directly in S3 without downloading objects.

Which solution will meet these requirements?

- A. Use Amazon Athena to query and filter the objects in Amazon S3.
- B. Use Amazon EMR to process and filter the objects.
- C. Use Amazon API Gateway to retrieve filtered results.
- D. Use Amazon ElastiCache to cache the objects.

#### ANSWER: A

##### Explanation:

Use Amazon Athena to query and filter the objects in Amazon S3. Athena is a serverless interactive query service that allows users to run standard SQL queries directly against data stored in Amazon S3. It is designed for data-lake analytics and can query common file formats such as CSV, JSON, Apache Parquet, and Apache ORC. Athena reads the relevant S3 data during query execution and returns only the requested results, so users do not need to download entire objects or manage servers and query infrastructure. Table metadata can be defined in the AWS Glue Data Catalog or another compatible metastore, allowing SQL queries to reference structured datasets stored in S3. For efficient filtering and lower scanning costs, datasets can be partitioned by commonly filtered fields and stored in columnar formats such as Parquet. Athena charges primarily for the amount of data scanned, making partitioning and optimized file formats useful practices for a data lake. AWS describes Athena as an analysis service that uses SQL to query data directly in Amazon S3. See the [Amazon Athena User Guide](#) and the [Amazon Athena product page](#).

#### QUESTION NO: 76

An ecommerce company is running a multi-tier application on AWS. The front-end and backend tiers run on Amazon EC2, and the database runs on Amazon RDS for MySQL. The backend tier communicates with the RDS instance. There are frequent calls to return identical database from the database that are causing performance slowdowns.

Which action should be taken to improve the performance of the backend?

- A. Implement Amazon SNS to store the database calls.
- B. Implement Amazon ElastiCache to cache frequently accessed database results.
- C. Implement an RDS for MySQL read replica to cache database calls.
- D. Implement Amazon Kinesis Data Firehose to stream the calls to the database.

#### ANSWER: B

##### Explanation:

Implement Amazon ElastiCache to cache frequently accessed database results. The workload repeatedly retrieves identical data, so an in-memory cache placed between the backend application and Amazon RDS for MySQL can serve subsequent requests without querying the database each time. ElastiCache supports Redis and Memcached, both of which provide very low-latency, high-throughput data retrieval for commonly accessed application data. The backend should first check the

cache for the requested result; on a cache miss, it retrieves the value from RDS and stores it in the cache with an appropriate expiration policy. This reduces read pressure, database CPU utilization, and connection contention on the RDS instance, while improving response times for customers. Cache invalidation or time-to-live settings should be selected to ensure cached results remain acceptably current when the underlying database data changes. AWS identifies database-query result caching as a common ElastiCache use case and recommends caching read-intensive, repeatedly accessed data to offload relational databases. See the [Amazon ElastiCache User Guide](#) and the [ElastiCache caching strategies documentation](#).

#### QUESTION NO: 77

A company wants to migrate from an on-premises data center to AWS. The data center hosts a storage server that stores data in an NFS-based file system. The storage server stores 200 GB of data. The company needs to migrate the data without interruption to existing services. Multiple resources in AWS must be able to access the data by using the NFS protocol.

Which combination of steps will meet these requirements MOST cost-effectively? (Select TWO.)

- A. Create an Amazon FSx for Lustre file system.
- B. Create an Amazon Elastic File System (Amazon EFS) file system.
- C. Create an Amazon S3 bucket to receive the data.
- D. Create an Amazon FSx for Windows file system.
- E. Install an AWS DataSync agent in the on-premises data center. Use a DataSync task between the on-premises file system and the AWS file system.

#### ANSWER: B E

##### Explanation:

Creating an Amazon Elastic File System (Amazon EFS) file system provides the required shared file storage destination in AWS. Amazon EFS is a fully managed, elastic file system that supports the NFS protocol and can be mounted concurrently by multiple compute resources, including Amazon EC2 instances, across multiple Availability Zones in a Region. This lets the migrated workload continue using standard file-system semantics rather than requiring application changes to use object storage or a different file-sharing protocol.

Installing an AWS DataSync agent in the on-premises data center and using a DataSync task between the on-premises file system and the AWS file system provides an efficient online migration path. DataSync supports NFS locations and Amazon EFS locations, transfers data securely, and can run repeated incremental synchronization tasks while the source remains active. The company can perform an initial copy, synchronize changed files as services continue operating, and then make a short final cutover after the remaining changes are copied. For a 200 GB dataset, this managed approach avoids deploying and operating custom transfer infrastructure while meeting the shared NFS access requirement. See [Amazon EFS User Guide](#) and [AWS DataSync User Guide](#).

#### QUESTION NO: 78

A company runs a web-based portal that provides users with global breaking news, local alerts, and weather updates. The portal delivers each user a personalized view by using mixture of static and dynamic content. Content is served over HTTPS through an API server running on an Amazon EC2 instance behind an Application Load Balancer (ALB). The company wants the portal to provide this content to its users across the world as quickly as possible.

How should a solutions architect design the application to ensure the LEAST amount of latency for all users?

- A. Deploy the application stack in a single AWS Region. Use Amazon CloudFront to serve all static and dynamic content by specifying the ALB as an origin.
- B. Deploy the application stack in two AWS Regions. Use an Amazon Route 53 latency routing policy to serve all content from the ALB in the closest Region.

C. Deploy the application stack in a single AWS Region. Use Amazon CloudFront to serve the static content. Serve the dynamic content directly from the ALB.

D. Deploy the application stack in two AWS Regions. Use an Amazon Route 53 geolocation routing policy to serve all content from the ALB in the closest Region.

**ANSWER: A**

**Explanation:**

Deploying the application stack in a single AWS Region and using Amazon CloudFront with the Application Load Balancer as the origin provides the lowest latency for a globally distributed audience. CloudFront receives viewer requests at an edge location close to the user, reducing the network distance and connection time required to reach the application. It can cache appropriate static assets, such as images, style sheets, and JavaScript, at edge locations for rapid delivery. CloudFront can also deliver personalized dynamic HTTPS requests to the ALB. For dynamic content that cannot be cached, CloudFront uses optimized connections between its edge network and the origin, which can improve performance compared with each user connecting directly to the Regional ALB across the public internet. Cache behaviors and headers can be configured so that personalized responses are not cached, while common content is cached safely. This design preserves the existing ALB-based application architecture while providing a global edge entry point for both types of content. AWS documents CloudFront as a service for accelerating both static and dynamic web content, including use of an Application Load Balancer as an origin. See [Amazon CloudFront Developer Guide](#) and [AWS guidance for dynamic content delivery with CloudFront](#).

**QUESTION NO: 79**

A city's weather forecast team is using Amazon DynamoDB in the data tier for an application. The application has several components. The analysis component of the application requires repeated reads against a large dataset. The application has started to temporarily consume all the read capacity in the DynamoDB table and is negatively affecting other applications that need to access the same data.

Which solution will resolve this issue with the LEAST development effort?

- A. Use DynamoDB Accelerator (DAX).
- B. Use Amazon CloudFront in front of DynamoDB.
- C. Create a DynamoDB table with a local secondary index (LSI).
- D. Use Amazon ElastiCache in front of DynamoDB.

**ANSWER: A**

**Explanation:**

Use DynamoDB Accelerator (DAX). DAX is a fully managed, DynamoDB-compatible, in-memory caching service designed specifically to reduce the read load placed on DynamoDB tables. It is well suited to an analysis component that repeatedly reads the same large dataset, because DAX can serve frequently requested data from its cache with microsecond latency rather than requiring each read to consume DynamoDB read capacity. This isolates the repetitive read workload from the table and helps preserve capacity for the other applications that must access the same data. DAX provides write-through caching and supports the DynamoDB API model, so applications generally make only minimal changes: they use the DAX client configuration instead of connecting directly to the DynamoDB endpoint. That compatibility and managed operation make it the solution requiring the least development effort while directly addressing temporary read-capacity exhaustion. AWS describes DAX as a fully managed, highly available in-memory cache for DynamoDB that can improve read performance by serving cached data without calls to the underlying table. See the [Amazon DynamoDB Developer Guide: DAX](#) and the [Amazon DynamoDB Accelerator product page](#).

**QUESTION NO: 80**

A media company uses an Amazon CloudFront distribution to deliver content over the internet. The company wants only premium customers to have access to the media streams and file content. The company stores all content in an Amazon S3

bucket. The company also delivers content on demand to customers for a specific purpose, such as movie rentals or music downloads.

Which solution will meet these requirements?

- A. Generate and provide S3 signed cookies to premium customers
- B. Generate and provide CloudFront signed URLs to premium customers.
- C. Use origin access control (OAC) to limit the access of non-premium customers
- D. Generate and activate field-level encryption to block non-premium customers.

**ANSWER: B**

**Explanation:**

“Generate and provide CloudFront signed URLs to premium customers.” is correct because CloudFront signed URLs provide restricted, time-limited access to individual objects delivered through a CloudFront distribution. After the company authenticates a premium customer through its application, it can generate a signed URL for the specific media stream or downloadable file that customer is entitled to access. CloudFront validates the URL signature before serving the object from its cache or retrieving it from the Amazon S3 origin.

A signed URL can include a policy that controls the expiration time and, when needed, other conditions such as the viewer IP address. This is well suited to on-demand entitlements: a movie rental URL can expire at the end of the rental period, and a music-download URL can be issued only for an authorized customer and a defined duration. The application retains control over who receives a URL, while CloudFront performs authorization enforcement at the edge. For private content, the S3 bucket should also be configured so CloudFront is the authorized path to the origin, preventing customers from bypassing the distribution with direct S3 access. AWS documents signed URLs as the mechanism for granting access to individual restricted files: [CloudFront signed URLs](#). See also [Serving private content through CloudFront](#).

## QUESTION NO: 81

A company has an application that uses an Amazon RDS for PostgreSQL database. The company is developing an application feature that will store sensitive information for an individual in the database.

During a security review of the environment, the company discovers that the RDS DB instance is not encrypting data at rest. The company needs a solution that will provide encryption at rest for all the existing data and for any new data that is entered for an individual.

Which combination of steps should the company take to meet these requirements? (Select TWO.)

- A. Create a snapshot of the DB instance. Enable encryption on the snapshot. Use the encrypted snapshot to create a new DB instance. Adjust the application configuration to use the new DB instance.
- B. Create a snapshot of the DB instance. Create an encrypted copy of the snapshot. Use the encrypted snapshot to create a new DB instance. Adjust the application configuration to use the new DB instance.
- C. Modify the configuration of the DB instance by enabling encryption. Create a snapshot of the DB instance. Use the snapshot to create a new DB instance. Adjust the application configuration to use the new DB instance.
- D. Use AWS Key Management Service (AWS KMS) to create a new default AWS managed aws/rds key. Select this key as the encryption key for operations with Amazon RDS.
- E. Use AWS Key Management Service (AWS KMS) to create a new customer managed key. Select this key as the encryption key for operations with Amazon RDS.

**ANSWER: B E**

**Explanation:**

Amazon RDS encryption at rest must be enabled when a DB instance is created; it cannot be enabled by modifying an existing unencrypted DB instance. To protect the existing PostgreSQL data, the company should first create a manual

snapshot of the current instance and then create an encrypted copy of that snapshot. The encrypted snapshot copy can be restored into a new RDS for PostgreSQL DB instance. Because the new instance is restored from an encrypted snapshot, its storage, automated backups, read replicas, and snapshots are encrypted at rest. After validating the restored database, the application can be updated to connect to the new DB instance so that both the migrated records and all subsequently written individual data are protected.

A customer managed AWS KMS key is appropriate as the encryption key for the snapshot-copy and restore operations. Customer managed keys give the company control over the key policy, grants, rotation configuration, and audit activity through AWS CloudTrail. This supports a security-sensitive workload while allowing RDS to use the key to encrypt and decrypt the database storage on the company's behalf. AWS documents the encrypted snapshot-copy and restore process at [Encrypting Amazon RDS resources](#) and describes customer managed KMS keys at [AWS KMS customer managed keys](#).

## QUESTION NO: 82

A company wants to grant an external vendor temporary, limited access to an Amazon S3 bucket to download files. The company does not want the external vendor to have access to the bucket for a long period of time.

Which solution will meet these requirements in the MOST secure way?

- A. Create an IAM user and programmatic access keys. Attach an IAM policy to the user that allows read-only access to the S3 bucket. Share the IAM user and programmatic access keys with the external vendor.
- B. Add a bucket policy to the S3 bucket that grants access based on the external vendor 's IP address range.
- C. Create a presigned URL for each required object in the S3 bucket. Share the presigned URLs with the external vendor.
- D. Create an IAM role and temporary access keys. Attach an IAM policy to the role that allows read-only access to the S3 bucket. Share the IAM role temporary access keys with the external vendor.

## ANSWER: C

### Explanation:

Create a presigned URL for each required object in the S3 bucket. Share the presigned URLs with the external vendor. This approach grants time-limited permission to perform a specific S3 operation, such as downloading an individual object, without giving the vendor AWS account credentials or general access to the bucket. The URL is signed using the permissions of the AWS principal that creates it, and the creator explicitly sets its expiration time. The vendor can use the URL with a browser or HTTP client until it expires, making it well suited to controlled external distribution of a known set of files.

Presigned URLs support least-privilege access because each URL can be scoped to one object and one intended operation. The company can therefore generate URLs only for the files the vendor needs and select a short validity period that matches the business requirement. For access keys created through IAM Identity Center or AWS STS credentials, the effective presigned-URL validity can also be limited by the lifetime of the underlying credentials. Treat presigned URLs as sensitive bearer tokens: anyone who possesses a valid URL can use it until expiration, so the company should transmit them through an appropriate secure channel. See the [Amazon S3 documentation on presigned URLs](#) and the [AWS guide to sharing objects with presigned URLs](#).

## QUESTION NO: 83

A company has an application that runs on Amazon EC2 instances and uses an Amazon Aurora database. The EC2 instances connect to the database by using user names and passwords that are stored locally in a file. The company wants to minimize the operational overhead of credential management.

What should a solutions architect do to accomplish this goal?

- A. Use AWS Secrets Manager. Turn on automatic rotation.
- B. Use AWS Systems Manager Parameter Store. Turn on automatic rotation.
- C. Create an Amazon S3 bucket to store objects that are encrypted with an AWS Key Management Service (AWS KMS) encryption key. Migrate the credential file to the S3 bucket. Point the application to the S3 bucket.

**D.** Create an encrypted Amazon Elastic Block Store (Amazon EBS) volume (or each EC2 instance. Attach the new EBS volume to each EC2 instance. Migrate the credential file to the new EBS volume. Point the application to the new EBS volume.

**ANSWER: A**

**Explanation:**

Use AWS Secrets Manager. Turn on automatic rotation. AWS Secrets Manager is designed to centrally store, retrieve, and manage sensitive values such as Amazon Aurora database usernames and passwords. The application running on Amazon EC2 can retrieve the secret at runtime by using an IAM role attached to the EC2 instances, rather than relying on a locally stored credential file. This removes the need to manually distribute, update, and protect credential files across individual instances.

Secrets Manager supports automatic rotation for Amazon RDS databases, including Aurora. Rotation updates the database credential and the corresponding stored secret on a defined schedule, helping the company maintain regularly changed credentials without building or operating its own rotation process. The application should retrieve the current secret from Secrets Manager, ideally using client-side caching where appropriate, so that it can continue to use the centrally managed credential after rotations occur. This approach reduces operational effort while improving credential security and auditability. See the [AWS Secrets Manager rotation documentation](#) and [retrieving secrets documentation](#).

**QUESTION NO: 84**

A company 's software development team needs an Amazon RDS Multi-AZ cluster. The RDS cluster will serve as a backend for a desktop client that is deployed on premises. The desktop client requires direct connectivity to the RDS cluster.

The company must give the development team the ability to connect to the cluster by using the client when the team is in the office.

Which solution provides the required connectivity MOST securely?

- A.** Create a VPC and two public subnets. Create the RDS cluster in the public subnets. Use AWS Site-to-Site VPN with a customer gateway in the company 's office.
- B.** Create a VPC and two private subnets. Create the RDS cluster in the private subnets. Use AWS Site-to-Site VPN with a customer gateway in the company 's office.
- C.** Create a VPC and two private subnets. Create the RDS cluster in the private subnets. Use RDS security groups to allow the company 's office IP ranges to access the cluster.
- D.** Create a VPC and two public subnets. Create the RDS cluster in the public subnets. Create a cluster user for each developer. Use RDS security groups to allow the users to access the cluster.

**ANSWER: B**

**Explanation:**

Create a VPC and two private subnets. Create the RDS cluster in the private subnets. Use AWS Site-to-Site VPN with a customer gateway in the company 's office. This design keeps the RDS Multi-AZ cluster on private network addresses, so the database is not exposed to direct connections from the public internet. A Site-to-Site VPN establishes encrypted IPsec tunnels between the company's on-premises network and the AWS VPC. Consequently, desktop clients connected to the office network can route privately to the cluster endpoint while employees are in the office, satisfying the requirement for direct client-to-database connectivity.

For the connection to function, the VPC route tables must include routes for the on-premises network through the VPN attachment, and the on-premises network must route the VPC CIDR through its customer gateway. The database security group should permit the required database port only from the relevant on-premises CIDR ranges (or a tightly controlled security-group-based source where applicable). Deploying the cluster across private subnets also aligns with the purpose of a DB subnet group for Multi-AZ RDS deployments: providing isolated subnets in multiple Availability Zones. AWS documents the encrypted connectivity and routing model for [AWS Site-to-Site VPN](#) and the networking considerations for [Amazon RDS in a VPC](#).

## QUESTION NO: 85

A company wants to use Amazon Elastic Container Service (Amazon ECS) to run its on-premises application in a hybrid environment. The application currently runs on containers on premises.

The company needs a single container solution that can scale in an on-premises, hybrid, or cloud environment. The company must run new application containers in the AWS Cloud and must use a load balancer for HTTP traffic.

Which combination of actions will meet these requirements? (Select TWO.)

- A.** Set up an ECS cluster that uses the AWS Fargate launch type for the cloud application containers. Use an Amazon ECS Anywhere external launch type for the on-premises application containers.
- B.** Set up an Application Load Balancer for cloud ECS services.
- C.** Set up a Network Load Balancer for cloud ECS services.
- D.** Set up an ECS cluster that uses the AWS Fargate launch type. Use Fargate for the cloud application containers and the on-premises application containers.
- E.** Set up an ECS cluster that uses the Amazon EC2 launch type for the cloud application containers. Use Amazon ECS Anywhere with an AWS Fargate launch type for the on-premises application containers.

## ANSWER: A B

### Explanation:

“Set up an ECS cluster that uses the AWS Fargate launch type for the cloud application containers. Use an Amazon ECS Anywhere external launch type for the on-premises application containers.” provides a single Amazon ECS control plane across the hybrid deployment. AWS Fargate runs the newly deployed cloud tasks without the company needing to provision or manage Amazon EC2 container instances. For the existing on-premises container environment, Amazon ECS Anywhere registers supported external instances with an ECS cluster and runs tasks by using the `EXTERNAL` launch type. This allows the company to use consistent ECS task definitions, services, scheduling, and operational tooling while placing workloads in the appropriate environment. AWS documents the ECS Anywhere model and external launch type at [Amazon ECS Anywhere](#).

“Set up an Application Load Balancer for cloud ECS services” fulfills the HTTP load-balancing requirement. An Application Load Balancer operates at Layer 7 and supports HTTP and HTTPS request routing. When used with an ECS service, it can distribute requests to healthy service tasks, integrate with target groups and health checks, and support dynamic task scaling as task capacity changes. This is the appropriate load-balancer model for HTTP application traffic in the AWS Cloud. See [Application Load Balancers](#) for its HTTP/HTTPS routing capabilities.

## QUESTION NO: 86

A hospital is designing a new application that gathers symptoms from patients. The hospital has decided to use Amazon Simple Queue Service (Amazon SQS) and Amazon Simple Notification Service (Amazon SNS) in the architecture.

A solutions architect is reviewing the infrastructure design. Data must be encrypted at rest and in transit. Only authorized personnel of the hospital should be able to access the data.

Which combination of steps should the solutions architect take to meet these requirements? (Select TWO.)

- A.** Turn on server-side encryption on the SQS components. Update the default key policy to restrict key usage to a set of authorized principals.
- B.** Turn on server-side encryption on the SNS components by using an AWS Key Management Service (AWS KMS) customer managed key. Apply a key policy to restrict key usage to a set of authorized principals.
- C.** Turn on encryption on the SNS components. Update the default key policy to restrict key usage to a set of authorized principals. Set a condition in the topic policy to allow only encrypted connections over TLS.

**D.** Turn on server-side encryption on the SQS components by using an AWS Key Management Service (AWS KMS) customer managed key. Apply a key policy to restrict key usage to a set of authorized principals. Set a condition in the queue policy to allow only encrypted connections over TLS.

**E.** Turn on server-side encryption on the SQS components by using an AWS Key Management Service (AWS KMS) customer managed key. Apply an IAM policy to restrict key usage to a set of authorized principals. Set a condition in the queue policy to allow only encrypted connections over TLS.

**ANSWER: B D**

**Explanation:**

Turn on server-side encryption on the SNS components by using an AWS Key Management Service (AWS KMS) customer managed key. Apply a key policy to restrict key usage to a set of authorized principals. This protects message content stored by Amazon SNS with envelope encryption, while a customer managed key gives the hospital direct control over the key policy. The policy can limit KMS cryptographic operations to the hospital's approved IAM principals and AWS services, supporting the requirement that only authorized personnel can access protected data.

Turn on server-side encryption on the SQS components by using an AWS Key Management Service (AWS KMS) customer managed key. Apply a key policy to restrict key usage to a set of authorized principals. Set a condition in the queue policy to allow only encrypted connections over TLS. This encrypts queued messages at rest and provides the necessary access control for KMS key use. The queue policy condition can require `aws:SecureTransport`, ensuring that requests to the queue use HTTPS/TLS and that data is protected while in transit. Together, these controls protect both messaging services at rest, enforce encrypted transport for queue access, and restrict decryption capabilities to authorized identities. See the [Amazon SQS server-side encryption documentation](#) and [Amazon SNS server-side encryption documentation](#).

**QUESTION NO: 87**

A company is developing an ecommerce application that will consist of a load-balanced front end, a container-based application, and a relational database. A solutions architect needs to create a highly available solution that operates with as little manual intervention as possible.

Which solutions meet these requirements? (Select TWO.)

- A.** Create an Amazon RDS DB instance in Multi-AZ mode.
- B.** Create an Amazon RDS DB instance and one or more replicas in another Availability Zone.
- C.** Create an Amazon EC2 instance-based Docker cluster to handle the dynamic application load.
- D.** Create an Amazon Elastic Container Service (Amazon ECS) cluster with a Fargate launch type to handle the dynamic application load.
- E.** Create an Amazon Elastic Container Service (Amazon ECS) cluster with an Amazon EC2 launch type to handle the dynamic application load.

**ANSWER: A D**

**Explanation:**

Create an Amazon RDS DB instance in Multi-AZ mode because Amazon RDS Multi-AZ deployments are designed for database high availability while minimizing operational work. RDS automatically provisions and maintains a synchronous standby database instance in a separate Availability Zone. If an infrastructure, Availability Zone, or database-instance failure occurs, RDS performs automatic failover and updates the database endpoint, allowing the application to reconnect without requiring the team to promote or manage a standby manually. This provides resilient relational storage for the ecommerce workload. See the [Amazon RDS Multi-AZ documentation](#).

Create an Amazon Elastic Container Service (Amazon ECS) cluster with a Fargate launch type to handle the dynamic application load because AWS Fargate is serverless compute for containers. ECS can maintain a desired number of application tasks, replace unhealthy tasks, distribute tasks across Availability Zones when configured with suitable subnets, and integrate with a load balancer. Fargate removes the need to provision, patch, maintain, or scale the EC2 instances that

would otherwise run the containers. ECS Service Auto Scaling can then adjust task count as demand changes, satisfying the requirement for high availability and minimal manual intervention. See [AWS Fargate for Amazon ECS](#).

### QUESTION NO: 88

A company generates approximately 20 GB of data multiple times each day. The company uses AWS DataSync to copy all data from on-premises storage to Amazon S3 every 6 hours for further processing. The analytics team wants to modify the copy process to copy only data relevant to the analytics team and ignore the rest of the data. The team wants to copy data as soon as possible and receive a notification when the copy process is finished. Which combination of steps will meet these requirements MOST cost-effectively? (Select THREE.)

- A.** Modify the data generation process on-premises to create a manifest file at the end of each data generation cycle with the names of the objects to be copied to Amazon S3. Create a custom script to upload the manifest file to an S3 bucket.
- B.** Modify the data generation process on-premises to create a manifest file at the end of the copy process with the names of the objects to be copied to Amazon S3. Create an AWS Lambda function to load the manifest file data into an Amazon DynamoDB table.
- C.** Create an AWS Lambda function that Amazon EventBridge invokes when the manifest file is loaded into Amazon DynamoDB. Configure the Lambda function to copy the data from on-premises storage to the S3 bucket that uses the manifest file.
- D.** Create an AWS Lambda function that an S3 Event Notification invokes when the manifest file is uploaded. Configure the Lambda function to invoke the DataSync task by calling the `StartTaskExecution` API action with a manifest.
- E.** Create an Amazon Simple Notification Service (Amazon SNS) topic. Create an Amazon EventBridge rule to send an email notification to the SNS topic when the DataSync task execution status changes to SUCCESS or to ERROR.
- F.** Create an Amazon Simple Notification Service (Amazon SNS) topic. Create an AWS Lambda function to send an email notification to the SNS topic when the DataSync task execution status changes to SUCCESS or to ERROR.

### ANSWER: A D E

#### Explanation:

Using a manifest-based AWS DataSync transfer is the cost-effective way to limit each execution to analytics-relevant files. The on-premises generation process can produce a manifest that lists the required source objects and upload that manifest to Amazon S3. DataSync supports manifests stored in S3, enabling a task execution to transfer only the files specified by the manifest rather than scanning and transferring the entire source dataset. This directly reduces unnecessary data movement and processing.

Having an S3 event invoke AWS Lambda when the manifest arrives provides event-driven execution. The function can call `StartTaskExecution` and provide the manifest configuration, so the transfer begins as soon as the relevant file list is available instead of waiting for a fixed six-hour schedule. This uses managed AWS services and avoids maintaining servers or a custom data-copy implementation.

Amazon EventBridge receives DataSync task-execution state-change events. An EventBridge rule can filter for terminal `SUCCESS` and `ERROR` states and publish to an Amazon SNS topic. SNS then delivers an email to subscribed analytics-team recipients, providing completion visibility for both successful and failed copies. See the [DataSync StartTaskExecution API documentation](#) and [DataSync EventBridge event documentation](#).

### QUESTION NO: 89

A company is storing data in Amazon S3 buckets. The company needs to retain any objects that contain personally identifiable information (PII) that might need to be reviewed.

A solutions architect must develop an automated solution to identify objects that contain PII and apply the necessary controls to prevent deletion before review.

Which combination of steps should the solutions architect take to meet these requirements? (Select THREE.)

- A. Create a job in Amazon Macie to scan the S3 buckets for the relevant sensitive data identifiers.
- B. Move the identified objects to the S3 Glacier Deep Archive storage class.
- C. Create an AWS Lambda function that performs an S3 Object Lock legal hold operation on the identified objects.
- D. Create an AWS Lambda function that applies an S3 Object Lock retention period to the identified objects in governance mode.
- E. Create an Amazon EventBridge rule that invokes the AWS Lambda function when Amazon Macie detects sensitive data.
- F. Configure multi-factor authentication (MFA) delete on the S3 buckets.

**ANSWER: A C E**

**Explanation:**

Creating a job in Amazon Macie to scan the S3 buckets for the relevant sensitive data identifiers provides the PII-discovery component. Macie sensitive-data discovery jobs analyze S3 objects using managed data identifiers and, when appropriate, custom data identifiers, and generate findings for objects containing sensitive data. Amazon Macie publishes finding events to Amazon EventBridge. Therefore, creating an Amazon EventBridge rule that invokes the AWS Lambda function when Amazon Macie detects sensitive data creates the event-driven automation needed to process each relevant finding.

The Lambda function can use the S3 `PutObjectLegalHold` operation to place a legal hold on each identified object version. A legal hold prevents the protected object version from being overwritten or deleted and remains in effect until it is explicitly removed after the review. This directly supports an indefinite, review-driven retention period rather than a predetermined expiration. The target bucket must have S3 Versioning and S3 Object Lock enabled; Object Lock is enabled when a bucket is created. See the [Amazon Macie documentation for monitoring finding events with EventBridge](#) and the [Amazon S3 Object Lock documentation](#).

**QUESTION NO: 90**

A company is running a critical workload on an Amazon RDS DB instance. The company needs the DB instance to be highly available. The company requires a recovery time of less than 5 minutes.

Which solution will meet these requirements?

- A. Create a read replica of the DB instance.
- B. Use AWS CloudFormation to create a template of the DB instance.
- C. Take periodic snapshots of the DB instance. Store the snapshots in Amazon S3.
- D. Modify the DB instance to use a Multi-AZ deployment.

**ANSWER: D**

**Explanation:**

Modify the DB instance to use a Multi-AZ deployment. Amazon RDS Multi-AZ deployments are specifically designed to provide high availability for production databases. RDS synchronously replicates data to a standby instance in a different Availability Zone within the same AWS Region. The standby is managed by RDS and is not intended to serve read traffic; its purpose is to provide resilient failover capability.

If the primary DB instance, its Availability Zone, or a supporting infrastructure component fails, Amazon RDS automatically fails over to the standby. The DB instance endpoint remains the same, so applications reconnect using the existing endpoint rather than requiring a manual DNS or connection-string update. This managed automatic failover minimizes operational intervention and is intended to restore database availability within a short recovery window, meeting a recovery-time objective of less than five minutes for this critical workload.

Multi-AZ can be enabled when creating an RDS DB instance or by modifying an existing instance. For implementation details and failover behavior, see the [Amazon RDS Multi-AZ deployments documentation](#) and the [Amazon RDS high availability documentation](#).

## QUESTION NO: 91

A company hosts a web application on Amazon EC2 instances behind an Application Load Balancer ALB. The company uses Amazon Route 53 to route traffic. The company also has a static website that is configured in an Amazon S3 bucket.

A solutions architect must use the static website as a backup to the web application. The failover to the static website must be fully automated.

Which combination of actions will meet these requirements? Select TWO.

- A. Create a primary failover routing policy record. Configure the value to be the ALB.
- B. Create an AWS Lambda function to switch from the primary website to the secondary website when the health check fails.
- C. Create a primary failover routing policy record. Configure the value to be the ALB. Associate the record with a Route 53 health check.
- D. Create a secondary failover routing policy record. Configure the value to be the static website. Associate the record with a Route 53 health check.
- E. Create a secondary failover routing policy record. Configure the value to be the static website.

**ANSWER: C E**

### Explanation:

Creating a primary failover routing policy record with the Application Load Balancer as its value and associating that record with a Route 53 health check provides the health-aware primary endpoint required for automated DNS failover. Route 53 monitors the configured health check and normally answers DNS queries with the primary record while the application is healthy. If the health check determines that the primary application is unhealthy, Route 53 stops returning that primary endpoint for failover queries.

Creating a secondary failover routing policy record with the static website as its value supplies the standby destination. Route 53 failover routing is specifically designed to pair one primary record with one secondary record of the same name and type. When the primary endpoint is unhealthy, Route 53 automatically returns the secondary record, directing new DNS resolutions to the Amazon S3 static website. This is a DNS-based, managed failover mechanism, so no custom orchestration is needed to change records during an outage. The secondary site should contain an appropriate maintenance page or a static version of the application so it can serve users immediately when Route 53 directs traffic to it. See [Amazon Route 53 failover routing](#) and [Configuring DNS failover in Route 53](#).

## QUESTION NO: 92

A company uses AWS to run its ecommerce platform. The platform is critical to the company's operations and has a high volume of traffic and transactions. The company configures a multi-factor authentication (MFA) device to secure its AWS account root user credentials. The company wants to ensure that it will not lose access to the root user account if the MFA device is lost.

Which solution will meet these requirements?

- A. Set up a backup administrator account that the company can use to log in if the company loses the MFA device.
- B. Add multiple MFA devices for the root user account to handle the disaster scenario.
- C. Create a new administrator account when the company cannot access the root account.
- D. Attach the administrator policy to another IAM user when the company cannot access the root account.

**ANSWER: B**

### Explanation:

Adding multiple MFA devices for the root user account to handle the disaster scenario is the appropriate solution because AWS permits multiple MFA devices to be registered for an AWS account root user. This provides resilient access without

removing the protection that MFA provides for the highly privileged root credentials. If a primary authenticator is lost, damaged, unavailable, or held by an unavailable authorized custodian, an authorized administrator can use another registered MFA device to complete the root-user sign-in process.

AWS documents that an AWS account root user can register up to eight MFA devices. The company should register independent devices and protect them using its security and break-glass access procedures—for example, storing a backup hardware authenticator securely and limiting its physical access to authorized personnel. This approach directly addresses the availability requirement while retaining MFA on the root user, which AWS recommends as a foundational account-security control. Root credentials should still be used only for the limited tasks that require root-user access, with day-to-day administration performed through appropriately controlled IAM identities.

See the AWS guidance on [MFA for the AWS account root user](#) and [enabling and managing MFA devices](#).

### QUESTION NO: 93

A company runs an application in two AWS Regions. The application stores customer profile data in Amazon DynamoDB. Users must read and update profile data with low latency in either Region. The company requires automatic multi-Region replication and must continue operating if one Region becomes unavailable.

Which combination of actions should a solutions architect take to meet these requirements? (Choose two.)

- A. Create a DynamoDB global table with a replica table in each Region.
- B. Configure the application in each Region to read from and write to the local DynamoDB replica.
- C. Create a DynamoDB table in one Region. Configure Amazon S3 Cross-Region Replication to replicate the table data.
- D. Create DynamoDB tables in both Regions. Use AWS Database Migration Service (AWS DMS) for ongoing replication.
- E. Create a DynamoDB table in each Region. Configure Amazon Route 53 failover routing to direct all writes to one table.

### ANSWER: A B

#### Explanation:

Creating a DynamoDB global table provides a fully managed, multi-Region, multi-active database. Each replica table can accept local reads and writes, and DynamoDB automatically replicates changes among the replica Regions. This supports continued application operation if one Region is unavailable.

Configuring the application in each Region to use its local global table replica minimizes latency because requests are served by DynamoDB in the same Region as the application. DynamoDB global tables are designed for multi-Region applications that require low-latency local access and disaster recovery capabilities. See the [Amazon DynamoDB global tables documentation](#).

### QUESTION NO: 94

A company is migrating a distributed application to AWS. The application serves variable workloads. The legacy platform consists of a primary server that coordinates jobs across multiple compute nodes. The company wants to modernize the application with a solution that maximizes resiliency and scalability.

How should a solutions architect design the architecture to meet these requirements?

- A. Configure an Amazon Simple Queue Service (Amazon SQS) queue as a destination for the jobs. Implement the compute nodes with Amazon EC2 instances that are managed in an Auto Scaling group. Configure EC2 Auto Scaling to use scheduled scaling.
- B. Configure an Amazon Simple Queue Service (Amazon SQS) queue as a destination for the jobs. Implement the compute nodes with Amazon EC2 instances that are managed in an Auto Scaling group. Configure EC2 Auto Scaling based on the size of the queue.

C. Implement the primary server and the compute nodes with Amazon EC2 instances that are managed in an Auto Scaling group. Configure AWS CloudTrail as a destination for the logs. Configure EC2 Auto Scaling based on the load on the primary server.

D. Implement the primary server and the compute nodes with Amazon EC2 instances that are managed in an Auto Scaling group. Configure Amazon EventBridge (Amazon CloudWatch Events) as a destination for the logs. Configure EC2 Auto Scaling based on the load on the compute nodes.

**ANSWER: B**

**Explanation:**

Configuring an Amazon Simple Queue Service (Amazon SQS) queue as a destination for the jobs, with compute nodes in an EC2 Auto Scaling group that scales according to queue size, decouples job submission from job processing. This removes the primary server as a coordination dependency and lets job producers continue to submit work even when individual worker instances fail, are replaced, or are being launched. Amazon SQS durably retains queued messages until workers successfully process and delete them, providing a resilient buffer during workload spikes and worker interruptions.

Queue-depth metrics, particularly the approximate number of visible messages, reflect the actual amount of unprocessed work. An EC2 Auto Scaling policy can use these Amazon CloudWatch metrics to add workers as the backlog grows and remove workers as it drains. A target-tracking design using backlog per instance is especially useful because it adjusts capacity in proportion to both the queue backlog and the number of active workers. This event-driven scaling approach accommodates unpredictable demand more effectively than a fixed schedule and allows independent scaling of the processing tier. See the AWS guidance for [scaling based on Amazon SQS queue metrics](#) and the [Amazon SQS Developer Guide](#).

**QUESTION NO: 95**

A company needs to store confidential files on AWS. The company accesses the files every week. The company must encrypt the files by using envelope encryption, and the encryption keys must be rotated automatically. The company must have an audit trail to monitor encryption key usage.

Which combination of solutions will meet these requirements? (Select TWO.)

- A. Store the confidential files in Amazon S3.
- B. Store the confidential files in Amazon S3 Glacier Deep Archive.
- C. Use server-side encryption with customer-provided keys (SSE-C).
- D. Use server-side encryption with Amazon S3 managed keys (SSE-S3).
- E. Use server-side encryption with AWS KMS managed keys (SSE-KMS).

**ANSWER: A E**

**Explanation:**

“Store the confidential files in Amazon S3.” provides durable object storage that is appropriate for files accessed weekly and integrates directly with server-side encryption using AWS Key Management Service. “Use server-side encryption with AWS KMS managed keys (SSE-KMS).” meets the encryption and key-management requirements. When Amazon S3 encrypts an object with SSE-KMS, it uses envelope encryption: S3 requests a unique data key from AWS KMS, uses that data key to encrypt the object data, and stores the encrypted data key alongside the object. The plaintext data key is not retained after encryption.

AWS KMS managed keys are automatically rotated by AWS KMS on an annual basis. AWS KMS also integrates with AWS CloudTrail, which records KMS API activity, including encryption-key usage such as GenerateDataKey, Encrypt, and Decrypt requests. This produces an auditable history that can identify the principal, time, source, and key involved in cryptographic operations. S3 with SSE-KMS therefore combines suitable frequent-access storage with managed envelope encryption, automatic rotation, and key-usage auditing. See the [Amazon S3 SSE-KMS documentation](#) and the [AWS KMS key rotation documentation](#).

## QUESTION NO: 96

A company is creating a prototype of an ecommerce website on AWS. The website consists of an Application Load Balancer, an Auto Scaling group of Amazon EC2 instances for web servers, and an Amazon RDS for MySQL DB instance that runs with the Single-AZ configuration.

The website is slow to respond during searches of the product catalog. The product catalog is a group of tables in the MySQL database that the company does not access frequently. A solutions architect has determined that the CPU utilization on the DB instance is high when product catalog searches occur.

What should the solutions architect recommend to improve the performance of the website during searches of the product catalog?

- A.** Migrate the product catalog to an Amazon Redshift database. Use the COPY command to load the product catalog tables.
- B.** Implement an Amazon ElastiCache for Redis cluster to cache the product catalog. Use lazy loading to populate the cache.
- C.** Add an additional scaling policy to the Auto Scaling group to launch additional EC2 instances when database response is slow.
- D.** Turn on the Multi-AZ configuration for the DB instance. Configure the EC2 instances to throttle the product catalog queries that are sent to the database.

**ANSWER: B**

### Explanation:

Implement an Amazon ElastiCache for Redis cluster to cache the product catalog. Use lazy loading to populate the cache. This is appropriate because product catalog information is read during searches but is not updated frequently, making it well suited to caching. The application can use a cache-aside (lazy-loading) pattern: it first checks Redis for the requested catalog data. On a cache hit, the application returns the data directly from the in-memory cache, avoiding a query to Amazon RDS for MySQL. On a cache miss, the application retrieves the data from MySQL, returns it to the user, and stores it in Redis for subsequent requests.

By serving repeated catalog searches from Redis, this design substantially reduces read-query demand and CPU utilization on the DB instance. Redis provides low-latency, in-memory access, so it also improves search response time for commonly requested products. Because catalog changes are infrequent, the application can choose an appropriate time-to-live value and invalidate or refresh affected cached entries when catalog data changes. AWS documents lazy loading as a caching strategy that loads data into the cache only when it is requested, which directly fits this workload. See the [Amazon ElastiCache caching strategies documentation](#) and the [Amazon ElastiCache for Redis documentation](#).

## QUESTION NO: 97

A company runs a NetApp storage array in an on-premises data center. The company wants to migrate the storage array to Amazon FSx for NetApp ONTAP. The company has a mix of NFS and SMB file shares with complex directory structures and over 60 million small files. The company has 10 Gbps of network bandwidth available. The company wants to optimize migration efficiency for the file system.

- A.** Use AWS DataSync with a bandwidth throttle. Use the All tiering policy.
- B.** Provision an AWS Storage Gateway Volume Gateway. Configure a zero-ETL integration with the FSx for NetApp ONTAP file system.
- C.** Set up NetApp SnapMirror replication between the on-premises array and the FSx for ONTAP file system.
- D.** Use AWS Snowball Edge to perform an offline migration.

**ANSWER: C**

### Explanation:

Set up NetApp SnapMirror replication between the on-premises array and the FSx for ONTAP file system. SnapMirror is purpose-built for replicating data between compatible NetApp ONTAP systems, including an on-premises ONTAP deployment and Amazon FSx for NetApp ONTAP. It transfers data at the storage level and performs an initial baseline transfer followed by incremental updates. This approach is particularly efficient for a very large namespace containing tens of millions of small files because it avoids a file-by-file migration workflow and preserves the ONTAP volume's data and relevant metadata during replication.

The company can create destination volumes on FSx for ONTAP, initialize SnapMirror relationships from the source volumes, and allow replication to use the available high-bandwidth network connection. Before cutover, it can run final incremental updates to reduce downtime. After the final transfer, the destination volume can be made writable and clients can be redirected to the FSx for ONTAP endpoints, which support both NFS and SMB access. This native NetApp replication workflow is designed to provide an efficient migration path while retaining the existing logical volume and directory organization. AWS documents SnapMirror as a supported method for migrating ONTAP data to FSx for ONTAP: [Migrating data to FSx for ONTAP](#) and [Working with SnapMirror](#).

## QUESTION NO: 98

A solutions architect is designing the cloud architecture for a new stateless application that will be deployed on AWS. The solutions architect created an Amazon Machine Image (AMI) and launch template for the application.

Based on the number of jobs that need to be processed, the processing must run in parallel while adding and removing application Amazon EC2 instances as needed. The application must be loosely coupled. The job items must be durably stored.

Which solution will meet these requirements?

- A.** Create an Amazon Simple Notification Service (Amazon SNS) topic to send the jobs that need to be processed. Create an Auto Scaling group by using the launch template with the scaling policy set to add and remove EC2 instances based on CPU usage.
- B.** Create an Amazon Simple Queue Service (Amazon SQS) queue to hold the jobs that need to be processed. Create an Auto Scaling group by using the launch template with the scaling policy set to add and remove EC2 instances based on network usage.
- C.** Create an Amazon Simple Queue Service (Amazon SQS) queue to hold the jobs that need to be processed. Create an Auto Scaling group by using the launch template with the scaling policy set to add and remove EC2 instances based on the number of items in the SQS queue.
- D.** Create an Amazon Simple Notification Service (Amazon SNS) topic to send the jobs that need to be processed. Create an Auto Scaling group by using the launch template with the scaling policy set to add and remove EC2 instances based on the number of messages published to the SNS topic.

## ANSWER: C

### Explanation:

Create an Amazon Simple Queue Service (Amazon SQS) queue to hold the jobs that need to be processed and create an Auto Scaling group that adds and removes EC2 instances based on the number of items in the SQS queue. Amazon SQS provides durable message storage and decouples job producers from the EC2-based workers that consume and process jobs. Each worker can retrieve jobs independently, which supports parallel processing and allows the stateless application instances to be added or removed without requiring direct coordination with job producers.

Amazon EC2 Auto Scaling can use Amazon CloudWatch metrics derived from SQS, such as the approximate number of visible messages, to adjust capacity according to the backlog. A backlog-aware scaling policy aligns the number of workers with the outstanding work: as queued jobs increase, the Auto Scaling group can launch more instances; as the backlog declines, it can reduce capacity. AWS also supports target tracking based on backlog per instance, a useful pattern for maintaining an appropriate queue-processing rate. For implementation details, see the [Amazon SQS basic architecture documentation](#) and [EC2 Auto Scaling guidance for scaling based on an SQS queue](#).

## QUESTION NO: 99

A company offers a food delivery service that is growing rapidly. Because of the growth, the company's order processing system is experiencing scaling problems during peak traffic hours. The current architecture includes the following:

- A group of Amazon EC2 instances that run in an Amazon EC2 Auto Scaling group to collect orders from the application
- Another group of EC2 instances that run in an Amazon EC2 Auto Scaling group to fulfill orders

The order collection process occurs quickly, but the order fulfillment process can take longer. Data must not be lost because of a scaling event.

A solutions architect must ensure that the order collection process and the order fulfillment process can both scale properly during peak traffic hours. The solution must optimize utilization of the company's AWS resources.

Which solution meets these requirements?

- A.** Use Amazon CloudWatch metrics to monitor the CPU of each instance in the Auto Scaling groups. Configure each Auto Scaling group's minimum capacity according to peak workload values.
- B.** Use Amazon CloudWatch metrics to monitor the CPU of each instance in the Auto Scaling groups. Configure a CloudWatch alarm to invoke an Amazon Simple Notification Service (Amazon SNS) topic that creates additional Auto Scaling groups on demand.
- C.** Provision two Amazon Simple Queue Service (Amazon SQS) queues: one for order collection and another for order fulfillment. Configure the EC2 instances to poll their respective queue. Scale the Auto Scaling groups based on notifications that the queues send.
- D.** Provision two Amazon Simple Queue Service (Amazon SQS) queues: one for order collection and another for order fulfillment. Configure the EC2 instances to poll their respective queue. Create a metric based on a backlog per instance calculation. Scale the Auto Scaling groups based on this metric.

### ANSWER: D

#### Explanation:

Provision two Amazon Simple Queue Service (Amazon SQS) queues: one for order collection and another for order fulfillment. Configure the EC2 instances to poll their respective queue. Create a metric based on a backlog per instance calculation. Scale the Auto Scaling groups based on this metric. This design decouples the two processing stages, so a surge of incoming orders can be durably buffered while fulfillment workers process messages at their own rate. Amazon SQS retains messages until they are successfully processed and deleted, which protects queued work when worker instances are terminated or when the fulfillment tier requires more time to complete an order.

A backlog-per-instance metric is well suited to queue-processing fleets because it connects scaling capacity directly to outstanding work. The metric can be calculated from the approximate number of visible messages in a queue divided by the number of in-service instances in the corresponding Auto Scaling group. Target tracking can then add instances as backlog rises and remove instances as it falls, allowing each tier to scale independently according to its processing speed and acceptable queue delay. This avoids maintaining peak capacity continuously and therefore improves resource utilization while preserving asynchronous, reliable order processing. AWS documents this queue-based scaling pattern and backlog-per-instance calculation in its [Amazon EC2 Auto Scaling SQS scaling guidance](#); SQS message retention and deletion behavior are described in the [Amazon SQS Developer Guide](#).

### QUESTION NO: 100

A company is migrating five on-premises applications to VPCs in the AWS Cloud. Each application is currently deployed in isolated virtual networks on premises and should be deployed similarly in the AWS Cloud. The applications need to reach a shared services VPC. All the applications must be able to communicate with each other.

If the migration is successful, the company will repeat the migration process for more than 100 applications.

Which solution will meet these requirements with the LEAST administrative overhead?

- A.** Deploy software VPN tunnels between the application VPCs and the shared services VPC. Add routes between the application VPCs in their subnets to the shared services VPC.

- B.** Deploy VPC peering connections between the application VPCs and the shared services VPC. Add routes between the application VPCs in their subnets to the shared services VPC through the peering connection.
- C.** Deploy an AWS Direct Connect connection between the application VPCs and the shared services VPC. Add routes from the application VPCs in their subnets to the shared services VPC and the applications VPCs. Add routes from the shared services VPC subnets to the applications VPCs.
- D.** Deploy a transit gateway with associations between the transit gateway and the application VPCs and the shared services VPC. Add routes between the application VPCs in their subnets and the application VPCs to the shared services VPC through the transit gateway.

**ANSWER: D**

**Explanation:**

Deploy a transit gateway with associations between the transit Gateway and the application VPCs and the shared services VPC. Add routes between the application VPCs and from the application VPCs to the shared services VPC through the transit gateway. This design uses AWS Transit Gateway as a central network-transit hub. Each isolated application VPC and the shared services VPC attaches to the transit gateway, allowing the organization to implement controlled, hub-and-spoke connectivity without creating and maintaining a growing mesh of individual network connections. Transit Gateway route tables determine which attachments can communicate, so the company can enable application-to-application connectivity and access to shared services while retaining separate VPCs for each application.

This architecture is designed to scale to large numbers of VPCs. Adding a migrated application generally consists of creating a VPC attachment, associating or propagating routes to the appropriate Transit Gateway route table, and configuring the VPC subnet route tables to use the attachment. This centralized routing model substantially reduces operational effort as the environment expands beyond 100 applications. AWS Transit Gateway supports thousands of VPC attachments per transit gateway, subject to AWS quotas, and supports route-table-based segmentation and traffic control. See the [AWS Transit Gateway documentation](#) and the [AWS Transit Gateway route tables documentation](#).

**QUESTION NO: 101**

A company has an application that generates a large number of files, each approximately 5 MB in size. The files are stored in Amazon S3. Company policy requires the files to be stored for 4 years before they can be deleted. Immediate accessibility is always required as the files contain critical business data that is not easy to reproduce. The files are frequently accessed in the first 30 days of the object creation but are rarely accessed after the first 30 days.

Which storage solution is MOST cost-effective?

- A.** Create an S3 bucket lifecycle policy to move files from S3 Standard to S3 Glacier 30 days from object creation. Delete the files 4 years after object creation.
- B.** Create an S3 bucket lifecycle policy to move files from S3 Standard to S3 One Zone-Infrequent Access (S3 One Zone-IA) 30 days from object creation. Delete the files 4 years after object creation.
- C.** Create an S3 bucket lifecycle policy to move files from S3 Standard to S3 Standard-Infrequent Access (S3 Standard-IA) 30 days from object creation. Delete the files 4 years after object creation.
- D.** Create an S3 bucket lifecycle policy to move files from S3 Standard to S3 Standard-Infrequent Access (S3 Standard-IA) 30 days from object creation. Move the files to S3 Glacier 4 years after object creation.

**ANSWER: C**

**Explanation:**

Create an S3 bucket lifecycle policy to move files from S3 Standard to S3 Standard-Infrequent Access (S3 Standard-IA) 30 days from object creation. Delete the files 4 years after object creation is the most cost-effective solution because it aligns storage cost with the known access pattern while preserving immediate availability and resilient storage for critical data. S3 Standard is appropriate during the first 30 days, when the objects are accessed frequently. After that period, S3 Standard-IA provides lower storage pricing for data that is accessed infrequently while retaining millisecond retrieval performance, so the files remain immediately accessible when needed. S3 Standard-IA also stores data redundantly across multiple Availability Zones, which is important for business-critical files that are difficult to reproduce. Each 5 MB object exceeds the 128 KB

minimum billable object size for this storage class, and transitioning after 30 days satisfies its minimum storage-duration consideration. An S3 Lifecycle expiration rule at four years automates enforcement of the retention requirement and removes objects once they are eligible for deletion. See the [Amazon S3 storage classes](#) overview and the [S3 Lifecycle transition considerations](#) documentation.

### QUESTION NO: 102

A company recently launched Linux-based application instances on Amazon EC2 in a private subnet and launched a Linux-based bastion host on an Amazon EC2 instance in a public subnet of a VPC. A solutions architect needs to connect from the on-premises network, through the company's internet connection to the bastion host and to the application servers. The solutions architect must make sure that the security groups of all the EC2 instances will allow that access.

Which combination of steps should the solutions architect take to meet these requirements? (Select TWO)

- A. Replace the current security group of the bastion host with one that only allows inbound access from the application instances
- B. Replace the current security group of the bastion host with one that only allows inbound access from the internal IP range for the company
- C. Replace the current security group of the bastion host with one that only allows inbound access from the external IP range for the company
- D. Replace the current security group of the application instances with one that allows inbound SSH access from only the private IP address of the bastion host
- E. Replace the current security group of the application instances with one that allows inbound SSH access from only the public IP address of the bastion host

### ANSWER: C D

#### Explanation:

The correct configuration allows administrators to establish an SSH connection from the company's on-premises network to the bastion host over the internet, then establish a separate SSH connection from that bastion host to each private application instance. "Replace the current security group of the bastion host with one that only allows inbound access from the external IP range for the company" correctly restricts inbound SSH to the public source CIDR range that AWS sees when on-premises users traverse the company internet connection. The bastion host is in a public subnet and is the only instance that needs to accept this internet-originated administrative traffic.

"Replace the current security group of the application instances with one that allows inbound SSH access from only the private IP address of the bastion host" correctly permits the second hop. For traffic sent from the bastion host to an application instance through the VPC, the source is the bastion host's private IP address, so the private instances can remain unreachable directly from the internet. Security groups are stateful, meaning response traffic for an allowed inbound SSH session is automatically permitted. AWS documents security-group rules and the use of CIDR source ranges at [Security group rules](#) and describes bastion hosts as controlled administrative entry points at [Controlling network access to AWS resources](#).

### QUESTION NO: 103

An ecommerce company hosts a three-tier web application in a VPC. The web tier runs on Amazon EC2 instances in two Availability Zones. The company stores a product catalog and customer sales information in Amazon DynamoDB.

The company's finance team uses a reporting application to generate reports of daily product sales. When the finance team runs the daily reports, a sudden performance decrease affects website customers.

The company wants to improve the performance of the system.

Which solution will meet these requirements with MINIMAL changes to the current architecture?

- A. Migrate the application to larger EC2 instances. Migrate the database to Amazon RDS for MySQL. Configure a read replica of the database in a second Availability Zone.

- B. Increase the compute capacity of the EC2 instances. Migrate the database to Amazon ElastiCache (Memcached).
- C. Implement DynamoDB Accelerator (DAX).
- D. Configure DynamoDB streams.

**ANSWER: C**

**Explanation:**

Implement DynamoDB Accelerator (DAX). DAX is a fully managed, in-memory cache designed specifically for Amazon DynamoDB. It can serve frequently accessed DynamoDB data with microsecond latency rather than requiring every applicable read to be processed directly by the DynamoDB table. By handling cacheable reads from the finance reporting workload, DAX reduces read pressure on the underlying tables, helping preserve responsive access for the customer-facing website during report generation.

This approach is particularly suitable because the application already uses DynamoDB and DAX is compatible with DynamoDB API operations. The required application change is generally limited to configuring the application to use a DAX cluster endpoint through the DAX client, rather than redesigning the data model or moving data to a different database engine. DAX provides write-through caching, so writes are sent to DynamoDB and cached data is updated accordingly. For reporting queries where eventually consistent reads are acceptable, DAX can substantially improve read performance while retaining DynamoDB as the system of record.

For implementation details and supported operations, see the [Amazon DynamoDB Accelerator documentation](#) and the [DAX consistency model documentation](#).

**QUESTION NO: 104**

A company is developing software that uses a PostgreSQL database schema. The company needs to configure development environments and test environments for its developers.

Each developer at the company uses their own development environment, which includes a PostgreSQL database. On average, each development environment is used for an 8-hour workday. The test environments will be used for load testing that can take up to 2 hours each day.

Which solution will meet these requirements MOST cost-effectively?

- A. Configure development environments and test environments with their own Amazon Aurora Serverless v2 PostgreSQL database.
- B. For each development environment, configure an Amazon RDS for PostgreSQL Single-AZ DB instance. For the test environment, configure a single Amazon RDS for PostgreSQL Multi-AZ DB instance.
- C. Configure development environments and test environments with their own Amazon Aurora PostgreSQL DB cluster.
- D. Configure an Amazon Aurora global database. Allow developers to connect to the database with their own credentials.

**ANSWER: A**

**Explanation:**

Configure development environments and test environments with their own Amazon Aurora Serverless v2 PostgreSQL database. Aurora Serverless v2 is a good fit for separate development and testing databases because it is compatible with PostgreSQL and automatically adjusts database capacity in response to workload demand. Each developer can have an isolated database environment without the company having to select and continuously provision a fixed database instance size for every environment.

This elasticity is particularly useful for the stated usage pattern: development databases are active only during a typical workday, while load-testing databases are used for short periods. Aurora Serverless v2 scales capacity in Aurora capacity units (ACUs) as demand changes, allowing lower capacity during idle or light-use periods and more capacity during load tests. Supported Aurora PostgreSQL versions can also be configured to auto-pause at zero ACUs, further reducing compute charges when an environment is inactive. Storage and certain other services are billed separately, but the on-demand

capacity model minimizes unnecessary provisioned database compute for intermittent workloads. See the [Amazon Aurora Serverless v2 documentation](#) and [Aurora Serverless v2 auto-pause documentation](#).

#### QUESTION NO: 105

A company runs critical workloads on Amazon EC2 instances. The company must create daily backups of the Amazon EBS volumes and retain the backups for 7 years. The backups must be protected from deletion or changes to the retention policy, including deletion by an administrator.

Which combination of actions should a solutions architect take to meet these requirements? (Choose two.)

- A. Create an AWS Backup plan that creates daily backups of the EC2 instances and retains the recovery points for 7 years.
- B. Configure AWS Backup Vault Lock in compliance mode for the backup vault.
- C. Create daily EBS snapshots manually. Copy the snapshots to a second AWS Region.
- D. Enable termination protection on the EC2 instances.
- E. Create an Amazon S3 lifecycle rule to retain the EBS volumes for 7 years.

#### ANSWER: A B

##### Explanation:

Create an AWS Backup plan with a daily backup schedule and a 7-year lifecycle retention period. AWS Backup provides centralized, policy-based backup management for supported AWS resources, including Amazon EC2 and Amazon EBS. A backup plan can assign resources by tags or resource IDs and automatically create recovery points according to the required schedule.

Configure AWS Backup Vault Lock in compliance mode on the backup vault. Compliance mode prevents users, including the AWS account root user, from deleting recovery points before the retention period expires or altering the Vault Lock configuration after the grace period. This provides the required immutable retention protection for the daily backups. See the [AWS Backup plan documentation](#) and the [AWS Backup Vault Lock documentation](#).

#### QUESTION NO: 106

A company is redesigning a static website. The company needs a solution to host the new website in the company 's AWS account. The solution must be secure and scalable.

Which combination of solutions will meet these requirements? (Select THREE.)

- A. Configure an Amazon CloudFront distribution. Set the Amazon S3 bucket as the origin.
- B. Associate an AWS Certificate Manager (ACM) TLS certificate to the Amazon CloudFront distribution.
- C. Enable static website hosting for the Amazon S3 bucket.
- D. Create an Amazon S3 bucket to store the static website content.
- E. Export the website 's SSL/TLS certificate from AWS Certificate Manager (ACM) to the root of the Amazon S3 bucket.
- F. Turn off Block Public Access for the Amazon S3 bucket.

#### ANSWER: A B D

##### Explanation:

**Configure an Amazon CloudFront distribution. Set the Amazon S3 bucket as the origin., Associate an AWS Certificate Manager (ACM) TLS certificate to the Amazon CloudFront distribution., and Create an Amazon S3 bucket to store the static website content.** form the appropriate architecture. Amazon S3 provides highly durable, scalable object

storage for the website's HTML, CSS, JavaScript, image, and other static assets. Amazon CloudFront serves as the public content delivery layer, caching content at edge locations to provide low latency and automatically scale for varying request volumes. CloudFront can use the S3 bucket origin through the standard S3 REST endpoint, which supports keeping the bucket private when access is restricted to CloudFront by using Origin Access Control.

An ACM certificate attached to the CloudFront distribution enables HTTPS for visitors using the site's custom domain name. For CloudFront, the ACM certificate must be requested or imported in the US East (N. Virginia) Region. Together, these services separate private content storage from global public delivery, provide encrypted client connections, and avoid the operational burden of managing web servers. AWS recommends using CloudFront Origin Access Control to securely permit CloudFront to retrieve objects from an S3 origin while retaining S3 Block Public Access settings. See [Amazon CloudFront documentation for Amazon S3 origins](#) and [restricting access to an S3 origin](#).

#### QUESTION NO: 107

A company uses AWS Organizations to create dedicated AWS accounts for each business unit to manage each business unit's account independently upon request. The root email recipient missed a notification that was sent to the root user email address of one account. The company wants to ensure that all future notifications are not missed. Future notifications must be limited to account administrators.

Which solution will meet these requirements?

- A.** Configure the company's email server to forward notification email messages that are sent to the AWS account root user email address to all users in the organization.
- B.** Configure all AWS account root user email addresses as distribution lists that go to a few administrators who can respond to alerts. Configure AWS account alternate contacts in the AWS Organizations console or programmatically.
- C.** Configure all AWS account root user email messages to be sent to one administrator who is responsible for monitoring alerts and forwarding those alerts to the appropriate groups.
- D.** Configure all existing AWS accounts and all newly created accounts to use the same root user email address. Configure AWS account alternate contacts in the AWS Organizations console or programmatically.

#### ANSWER: B

##### Explanation:

Configure all AWS account root user email addresses as distribution lists that go to a few administrators who can respond to alerts. Configure AWS account alternate contacts in the AWS Organizations console or programmatically. This approach ensures that messages delivered to each account's root email address are received by multiple designated account administrators rather than relying on a single individual. Using a distribution list also allows the company to maintain administrator membership centrally, so access can be updated as responsibilities change without changing the AWS account root credentials or email address.

Alternate contacts provide dedicated recipients for billing, operations, and security communications. AWS Organizations can centrally manage these contacts for member accounts when trusted access for AWS account management is enabled, either through the Organizations console or the AWS Account Management API. This provides an account-specific, scalable notification configuration while keeping delivery limited to the small group of administrators responsible for each account. AWS documents alternate contacts and their notification categories at [Update alternate contacts for an AWS account](#). Central management across an organization is documented at [Updating alternate contacts for member accounts](#).

#### QUESTION NO: 108

A company uses Amazon EC2 instances behind an Application Load Balancer (ALB) to serve content to users. The company uses Amazon Elastic Block Store (Amazon EBS) volumes to store data.

The company needs to encrypt data in transit and at rest.

Which combination of services will meet these requirements? (Select TWO.)

- A.** Amazon GuardDuty

- B. AWS Shield
- C. AWS Certificate Manager (ACM)
- D. AWS Secrets Manager
- E. AWS Key Management Service (AWS KMS)

**ANSWER: C E**

**Explanation:**

**AWS Certificate Manager (ACM)** meets the data-in-transit requirement by provisioning and managing SSL/TLS certificates that can be attached to an Application Load Balancer HTTPS listener. The ALB uses that certificate to terminate TLS connections from users, encrypting traffic as it travels between clients and the load balancer. ACM also handles certificate deployment and managed renewal for eligible ACM-issued public certificates, reducing certificate-management overhead. If encryption is also required from the ALB to the EC2 targets, the target group can use HTTPS, with the application instances configured for TLS.

**AWS Key Management Service (AWS KMS)** meets the data-at-rest requirement for Amazon EBS. EBS encryption encrypts the volume data, disk I/O, snapshots created from the volume, and volumes created from those snapshots. When creating or enabling encryption for an EBS volume, the company can select an AWS KMS key, including the AWS managed key for EBS or a customer managed KMS key. KMS provides centralized key control, authorization through key policies and IAM, and auditing of key use through AWS CloudTrail.

References: [Application Load Balancer HTTPS listeners and ACM certificates](#); [Amazon EBS encryption](#).

**QUESTION NO: 109**

A company has an organization in AWS Organizations. The company runs Amazon EC2 instances across four AWS accounts in the root organizational unit (OU). There are three nonproduction accounts and one production account. The company wants to prohibit users from launching EC2 instances of a certain size in the nonproduction accounts. The company has created a service control policy (SCP) to deny access to launch instances that use the prohibited types.

Which solutions to deploy the SCP will meet these requirements? (Select TWO.)

- A. Attach the SCP to the root OU for the organization.
- B. Attach the SCP to the three nonproduction Organizations member accounts.
- C. Attach the SCP to the Organizations management account.
- D. Create an OU for the production account. Attach the SCP to the OU. Move the production member account into the new OU.
- E. Create an OU for the required accounts. Attach the SCP to the OU. Move the nonproduction member accounts into the new OU.

**ANSWER: B E**

**Explanation:**

Attaching the SCP to the three nonproduction Organizations member accounts meets the requirement because SCPs can be attached directly to individual member accounts. The explicit deny on the specified EC2 instance types then establishes a maximum permission boundary for principals in each of those nonproduction accounts. The production account remains outside the SCP attachment scope, so its users can continue to launch the otherwise prohibited instance sizes, subject to its own IAM permissions and any other applicable policies.

Creating an OU for the required accounts, attaching the SCP to the OU, and moving the nonproduction member accounts into the new OU also meets the requirement. SCPs attached to an OU are inherited by every account in that OU, making this approach an effective way to centrally apply the same preventive control across all current nonproduction accounts. It also supports scalable governance: additional nonproduction accounts can receive the restriction simply by placing them in that OU. SCPs define the maximum available permissions and apply to member accounts in the policy attachment hierarchy; an

explicit deny prevents the restricted EC2 launch requests even when an IAM policy would otherwise allow them. AWS documents SCP attachment and inheritance behavior at [Service control policies \(SCPs\)](#) and describes OU-based account organization at [Managing organizational units](#).

#### QUESTION NO: 110

A company uses an Amazon CloudFront distribution to serve content pages for its website. The company needs to ensure that clients use a TLS certificate when accessing the company's website. The company wants to automate the creation and renewal of the TLS certificates.

Which solution will meet these requirements with the MOST operational efficiency?

- A. Use a CloudFront security policy to create a certificate.
- B. Use a CloudFront origin access control (OAC) to create a certificate.
- C. Use AWS Certificate Manager (ACM) to create a certificate. Use DNS validation for the domain.
- D. Use AWS Certificate Manager (ACM) to create a certificate. Use email validation for the domain.

#### ANSWER: C

##### Explanation:

Use AWS Certificate Manager (ACM) to create a certificate. Use DNS validation for the domain. ACM integrates directly with Amazon CloudFront to provide TLS certificates for viewer connections over HTTPS. For a CloudFront distribution, the ACM certificate must be requested or imported in the US East (N. Virginia) Region, `us-east-1`, and then associated with the distribution's alternate domain name.

DNS validation provides the greatest operational efficiency because ACM supplies one or more CNAME validation records that can be added to the domain's DNS configuration. After ACM detects the records and issues the certificate, it can automatically renew an ACM-issued public certificate as long as the validation CNAME records remain in DNS and the certificate continues to be in use. This avoids recurring manual approval activities and supports uninterrupted HTTPS service for the CloudFront hostname. The distribution can also be configured to redirect HTTP viewers to HTTPS, ensuring clients access the site through TLS.

AWS documents CloudFront's ACM regional requirement and certificate association process in the [CloudFront alternate domain name and HTTPS requirements](#). ACM's DNS validation and managed renewal behavior are described in the [ACM DNS validation documentation](#).

#### QUESTION NO: 111

A company runs multiple web applications on Amazon EC2 instances behind a single Application Load Balancer (ALB). The application experiences unpredictable traffic spikes throughout each day. The traffic spikes cause high latency. The unpredictable spikes last less than 3 hours. The company needs a solution to resolve the latency issue caused by traffic spikes.

- A. Use EC2 instances in an Auto Scaling group. Configure the ALB and Auto Scaling group to use a target tracking scaling policy.
- B. Use EC2 Reserved Instances in an Auto Scaling group. Configure the Auto Scaling group to use a scheduled scaling policy based on peak traffic hours.
- C. Use EC2 Spot Instances in an Auto Scaling group. Configure the Auto Scaling group to use a scheduled scaling policy based on peak traffic hours.
- D. Use EC2 Reserved Instances in an Auto Scaling group. Replace the ALB with a Network Load Balancer (NLB).

#### ANSWER: A

**Explanation:**

Use EC2 instances in an Auto Scaling group. Configure the ALB and Auto Scaling group to use a target tracking scaling policy. This approach is designed for variable and unpredictable demand because target tracking continuously monitors a chosen CloudWatch metric and adjusts the number of healthy EC2 instances to maintain a specified target value. For an ALB-based web application, an especially relevant metric is `ALBRequestCountPerTarget`; alternatively, average CPU utilization can be used when it accurately represents application load. As traffic rises, the Auto Scaling group adds instances and the ALB distributes requests across the expanded target group, reducing per-instance load and latency. When the short-lived spike subsides, Auto Scaling removes unneeded instances, so the company avoids continuously operating excess capacity. Target tracking policies are dynamic rather than dependent on knowing the time of a future surge, making them appropriate for spikes that occur at unpredictable times. The Auto Scaling group should have suitable minimum, desired, and maximum capacity values, and its instances should be registered with the ALB target group. AWS documents target tracking as a policy type that increases or decreases capacity based on a metric and target value. See the [Amazon EC2 Auto Scaling target tracking documentation](#) and the [dynamic scaling policy documentation](#).

**QUESTION NO: 112**

An ecommerce company runs several internal applications in multiple AWS accounts. The company uses AWS Organizations to manage its AWS accounts.

A security appliance in the company's networking account must inspect interactions between applications across AWS accounts.

Which solution will meet these requirements?

- A. Deploy a Network Load Balancer (NLB) in the networking account to send traffic to the security appliance. Configure the application accounts to send traffic to the NLB by using an interface VPC endpoint in the application accounts
- B. Deploy an Application Load Balancer (ALB) in the application accounts to send traffic directly to the security appliance.
- C. Deploy a Gateway Load Balancer (GWLB) in the networking account to send traffic to the security appliance. Configure the application accounts to send traffic to the GWLB by using Gateway Load Balancer endpoints in the application VPCs.
- D. Deploy an interface VPC endpoint in the application accounts to send traffic directly to the security appliance.

**ANSWER: C****Explanation:**

Deploying a Gateway Load Balancer (GWLB) in the networking account and Gateway Load Balancer endpoints in the application VPCs provides a centralized, scalable traffic-inspection architecture. GWLB is purpose-built to insert virtual network appliances, such as firewalls, intrusion-detection systems, and deep-packet-inspection appliances, transparently into a traffic path. The GWLB distributes flows to the security appliance fleet while preserving flow symmetry, so packets in the same connection are consistently handled by the appropriate appliance.

Each application VPC uses a Gateway Load Balancer endpoint as the route target for traffic that requires inspection. The endpoint securely forwards encapsulated traffic to the GWLB service in the networking account, where the GWLB sends it to the registered appliance targets. This hub-and-spoke design allows the organization to operate and scale security appliances centrally instead of deploying and managing a separate appliance stack in every application account. Appropriate VPC route tables direct the relevant inter-application traffic through the endpoint, enabling inspection before the traffic reaches its destination. AWS documents GWLB endpoints as the mechanism for connecting VPC traffic to GWLB services and supports centralized inspection designs across VPCs and accounts. See [Gateway Load Balancer documentation](#) and [Gateway Load Balancer endpoint documentation](#).

**QUESTION NO: 113**

A company uses AWS Organizations to manage hundreds of AWS accounts. The company must centrally record API activity from all current and future member accounts. The logs must be stored in a dedicated log archive account and must not be modifiable by administrators in the member accounts.

Which combination of actions should a solutions architect take to meet these requirements? (Choose two.)

- A. Create an organization trail in AWS CloudTrail.
- B. Configure the organization trail to deliver logs to an Amazon S3 bucket in the log archive account.
- C. Create an individual CloudTrail trail manually in each member account.
- D. Configure Amazon CloudWatch Logs agents on all Amazon EC2 instances in every account.
- E. Create an Amazon S3 bucket in each member account. Use S3 Cross-Region Replication to copy the logs to the log archive account.

**ANSWER: A B**

**Explanation:**

Creating an organization trail in AWS CloudTrail records events for all accounts in the organization, including accounts that are created or added later. An organization trail is managed from the management account or a designated delegated administrator account and provides centralized audit coverage.

Configuring the organization trail to deliver logs to an S3 bucket in a dedicated log archive account separates log storage from member-account administration. The bucket policy can allow CloudTrail to write logs while preventing member-account administrators from changing or deleting the logs. See the [AWS CloudTrail organization trails documentation](#).

**QUESTION NO: 114**

A global company is using Amazon API Gateway to design REST APIs for its loyalty club users in the us-east-1 Region and the ap-southeast-2 Region. A solutions architect must design a solution to protect these API Gateway managed REST APIs across multiple accounts from SQL injection and cross-site scripting attacks.

Which solution will meet these requirements with the LEAST amount of administrative effort?

- A. Set up AWS WAF in both Regions. Associate Regional web ACLs with an API stage.
- B. Set up AWS Firewall Manager in both Regions. Centrally configure AWS WAF rules.
- C. Set up AWS Shield in both Regions. Associate Regional web ACLs with an API stage.
- D. Set up AWS Shield in one of the Regions. Associate Regional web ACLs with an API stage.

**ANSWER: B**

**Explanation:**

Set up AWS Firewall Manager in both Regions. Centrally configure AWS WAF rules is the correct solution because AWS Firewall Manager provides centralized security-policy administration across an AWS Organizations environment. A security team can create a Firewall Manager AWS WAF policy that applies a standardized web ACL configuration to eligible API Gateway REST API stages in multiple member accounts. The policy can include AWS WAF protections for SQL injection and cross-site scripting, such as SQLi and XSS match statements or AWS Managed Rules, and Firewall Manager automatically creates and maintains the required web ACL associations for resources that fall within the policy scope.

This approach is particularly suited to a multi-account architecture because it avoids requiring teams to manually create, configure, and associate equivalent regional web ACLs in every account. The policy can be scoped to the required Regions, including us-east-1 and ap-southeast-2, while retaining centralized visibility and ongoing enforcement as APIs are added or changed. AWS WAF is the service that inspects HTTP(S) requests and supports rules that detect SQL injection and cross-site scripting patterns; Firewall Manager reduces the operational overhead of deploying those protections consistently at organizational scale. See the [AWS Firewall Manager documentation for AWS WAF policies](#) and the [AWS WAF SQL injection match rule documentation](#).

**QUESTION NO: 115**

A company runs an order management application on AWS. The application allows customers to place orders and pay with a credit card. The company uses an Amazon CloudFront distribution to deliver the application.

A security team has set up logging for all incoming requests. The security team needs a solution to generate an alert if any user modifies the logging configuration.

Options (Select TWO):

- A.** Configure an Amazon EventBridge rule that is invoked when a user creates or modifies a CloudFront distribution. Add the AWS Lambda function as a target of the EventBridge rule.
- B.** Create an Application Load Balancer (ALB). Enable AWS WAF rules for the ALB. Configure an AWS Config rule to detect security violations.
- C.** Create an AWS Lambda function to detect changes in CloudFront distribution logging. Configure the Lambda function to use Amazon Simple Notification Service (Amazon SNS) to send notifications to the security team.
- D.** Set up Amazon GuardDuty. Configure GuardDuty to monitor findings from the CloudFront distribution. Create an AWS Lambda function to address the findings.
- E.** Create a private API in Amazon API Gateway. Use AWS WAF rules to protect the private API from common security problems.

**ANSWER: A C**

**Explanation:**

Configuring an Amazon EventBridge rule that is invoked when a user creates or modifies a CloudFront distribution and adding an AWS Lambda function as its target provides the event-driven detection mechanism. CloudFront management API activity is recorded by AWS CloudTrail, and CloudTrail events can be delivered to Amazon EventBridge. An EventBridge rule can filter for CloudFront distribution management events, including distribution update activity, so the response begins promptly when a configuration change occurs. The Lambda function can inspect the CloudTrail event details, including the requested distribution configuration, to determine whether the logging settings were changed rather than alerting indiscriminately for every distribution update.

Creating an AWS Lambda function to detect changes in CloudFront distribution logging and configuring it to use Amazon Simple Notification Service (Amazon SNS) to send notifications to the security team provides the alerting portion of the solution. After EventBridge invokes the function, the function can publish a message to an SNS topic when it identifies a logging-configuration modification. SNS then distributes the alert to subscribed security-team endpoints, such as email, SMS, HTTPS endpoints, or an incident-management integration. This combination separates detection, validation, and notification while remaining serverless and near real time. See the [CloudFront logging and monitoring documentation](#) and [Amazon EventBridge integration with CloudTrail events](#).

## QUESTION NO: 116

A company decides to use AWS Key Management Service (AWS KMS) for data encryption operations. The company must create a KMS key and automate the rotation of the key. The company also needs the ability to deactivate the key and schedule the key for deletion.

Which solution will meet these requirements?

- A.** Create an asymmetric customer managed KMS key. Enable automatic key rotation.
- B.** Create a symmetric customer managed KMS key. Disable the envelope encryption option.
- C.** Create a symmetric customer managed KMS key. Enable automatic key rotation.
- D.** Create an asymmetric customer managed KMS key. Disable the envelope encryption option.

**ANSWER: C**

**Explanation:**

Create a symmetric customer managed KMS key and enable automatic key rotation. Customer managed KMS keys provide the key lifecycle controls required here: an authorized principal can disable a key, which prevents cryptographic operations using that key, and can schedule the key for deletion. AWS KMS applies a mandatory waiting period when deletion is scheduled, allowing the organization to cancel deletion if necessary. These controls are available because the company manages the key rather than AWS managing it on the company's behalf.

A symmetric customer managed KMS key used for encryption and decryption also supports AWS KMS automatic key rotation. When enabled, AWS KMS creates new backing key material on its rotation schedule while retaining the same KMS key ARN, key ID, key policy, and permissions. Applications can therefore continue referring to the same KMS key without changes while newly generated ciphertext uses the current key material and AWS KMS preserves prior material for decryption of existing ciphertext. This directly combines automated rotation with the required ability to deactivate and retire the key through scheduled deletion. AWS documents rotation eligibility and behavior at [Rotating KMS keys](#), and documents disabling and deletion lifecycle operations at [Deleting KMS keys](#).

#### QUESTION NO: 117

A company runs database workloads on AWS that are the backend for the company 's customer portals. The company runs a Multi-AZ database cluster on Amazon RDS for PostgreSQL.

The company needs to implement a 30-day backup retention policy. The company currently has both automated RDS backups and manual RDS backups. The company wants to maintain both types of existing RDS backups that are less than 30 days old.

Which solution will meet these requirements MOST cost-effectively?

- A.** Configure the RDS backup retention policy to 30 days for automated backups by using AWS Backup. Manually delete manual backups that are older than 30 days.
- B.** Disable RDS automated backups. Delete automated backups and manual backups that are older than 30 days. Configure the RDS backup retention policy to 30 days for automated backups.
- C.** Configure the RDS backup retention policy to 30 days for automated backups. Manually delete manual backups that are older than 30 days
- D.** Disable RDS automated backups. Delete automated backups and manual backups that are older than 30 days automatically by using AWS CloudFormation. Configure the RDS backup retention policy to 30 days for automated backups.

#### ANSWER: C

#### Explanation:

Configure the RDS backup retention policy to 30 days for automated backups. Manually delete manual backups that are older than 30 days is the most cost-effective solution because Amazon RDS manages automated backup retention natively. Setting the DB cluster's backup retention period to 30 days preserves automated backups within that window and causes RDS to remove automated recovery data after the retention period expires. This supports point-in-time recovery throughout the configured retention period without requiring a separate backup-management service.

Manual RDS backups are DB snapshots created explicitly by a user or process. They are retained until they are deleted, rather than being governed by the automated-backup retention setting. Deleting only the manual snapshots older than 30 days preserves all existing manual snapshots that are still within the required period while preventing indefinite snapshot-storage charges. This approach uses the built-in RDS retention control for automated backups and targeted cleanup for manual snapshots, meeting the retention requirement without introducing unnecessary automation infrastructure or additional services.

For details, see the [Amazon RDS automated backup retention documentation](#) and the [Amazon RDS backup and snapshot documentation](#).

#### QUESTION NO: 118

A company is hosting multiple websites for several lines of business under its registered parent domain. Users accessing these websites will be routed to appropriate backend Amazon EC2 instances based on the subdomain. The websites host static webpages, images, and server-side scripts like PHP and JavaScript.

Some of the websites experience peak access during the first two hours of business with constant usage throughout the rest of the day. A solutions architect needs to design a solution that will automatically adjust capacity to these traffic patterns while keeping costs low.

Which combination of AWS services or features will meet these requirements? (Select TWO.)

- A. AWS Batch
- B. Network Load Balancer
- C. Application Load Balancer
- D. Amazon EC2 Auto Scaling
- E. Amazon S3 website hosting

**ANSWER: C D**

**Explanation:**

**Application Load Balancer** is appropriate because it operates at the application layer and supports host-based routing. The load balancer listener rules can inspect the HTTP `Host` header and forward requests for each subdomain, such as `sales.example.com` or `support.example.com`, to the corresponding target group of Amazon EC2 instances. This allows the company to use one load-balancing entry point while directing each line of business to its own backend application fleet. Application Load Balancers also support HTTP and HTTPS workloads, making them suitable for serving dynamic PHP applications as well as static web content. AWS documents host-header conditions and routing rules at [Application Load Balancer listener rules](#).

**Amazon EC2 Auto Scaling** provides the capacity adjustment mechanism. An Auto Scaling group can increase the number of healthy EC2 instances during the recurring morning demand period and reduce capacity afterward, maintaining availability while avoiding the cost of running peak-level capacity throughout the day. Because the traffic pattern is predictable, scheduled scaling can set the desired capacity ahead of the first two business hours; dynamic scaling policies can also respond to metrics such as average CPU utilization or request count per target. This combination provides both subdomain-aware request routing and cost-efficient, automatic compute scaling. See [Amazon EC2 Auto Scaling](#) for scaling policy and scheduled scaling guidance.

#### QUESTION NO: 119

A company recently migrated its web application to AWS by rehosting the application on Amazon EC2 instances in a single AWS Region. The company wants to redesign its application architecture to be highly available and fault tolerant. Traffic must reach all running EC2 instances randomly.

Which combination of steps should the company take to meet these requirements? (Choose two.)

- A. Create an Amazon Route 53 failover routing policy.
- B. Create an Amazon Route 53 weighted routing policy.
- C. Create an Amazon Route 53 multivalue answer routing policy.
- D. Launch three EC2 instances: two instances in one Availability Zone and one instance in another Availability Zone.
- E. Launch four EC2 instances: two instances in one Availability Zone and two instances in another Availability Zone.

**ANSWER: C E**

**Explanation:**

**Create an Amazon Route 53 multivalue answer routing policy and Launch four EC2 instances: two instances in one Availability Zone and two instances in another Availability Zone.** together provide DNS-level traffic distribution and infrastructure redundancy within the Region. Route 53 multivalue answer routing returns multiple healthy resource records in response to DNS queries. It selects these records randomly, allowing clients to receive addresses for different healthy EC2 instances rather than consistently targeting only one instance. Route 53 health checks can be associated with the records, so instances that fail health evaluation are not returned in responses while healthy instances remain available.

Deploying instances across two Availability Zones protects the application from an Availability Zone failure. Placing two instances in each zone preserves capacity and redundancy even if one instance fails or an entire zone becomes unavailable. This design therefore provides multiple independently reachable application instances and distributes DNS answers among the healthy instances. The architecture meets the stated random-routing requirement without requiring a centralized load balancer. For details, see the [Amazon Route 53 multivalue answer routing documentation](#) and the [AWS Regions and Availability Zones overview](#).

#### QUESTION NO: 120

A manufacturing company runs an order processing application in its VPC. The company wants to securely send messages from the application to an external Salesforce system that uses Open Authorization (OAuth).

A solutions architect needs to integrate the company 's order processing application with the external Salesforce system.

Which solution will meet these requirements?

- A.** Create an Amazon Simple Notification Service (Amazon SNS) topic in a fanout configuration that pushes data to an HTTPS endpoint. Configure the order processing application to publish messages to the SNS topic.
- B.** Create an Amazon Simple Notification Service (Amazon SNS) topic in a fanout configuration that pushes data to an Amazon Data Firehose delivery stream that has a HTTP destination. Configure the order processing application to publish messages to the SNS topic.
- C.** Create an Amazon EventBridge API destination with an OAuth connection and an EventBridge rule. Configure the order processing application to publish messages to Amazon EventBridge.
- D.** Create an Amazon Managed Streaming for Apache Kafka (Amazon MSK) topic that has an outbound MSK Connect connector. Configure the order processing application to publish messages to the MSK topic.

#### ANSWER: C

##### Explanation:

Create an Amazon EventBridge API destination with an EventBridge connection configured for OAuth, then use an EventBridge rule to route the application's events to that destination. The order processing application can publish custom events to an EventBridge event bus by using the PutEvents API. The rule can filter or transform those events as needed before forwarding them to the Salesforce HTTPS endpoint.

API destinations are designed specifically for sending EventBridge events to external HTTP API endpoints. An EventBridge connection securely stores the authorization configuration, including OAuth client credentials, and EventBridge manages obtaining and refreshing OAuth access tokens for requests to the destination. This avoids embedding Salesforce credentials or token-management logic in the order processing application. API destinations also provide managed delivery behavior, including configurable rate limits and retry handling, which supports a resilient integration with the external service.

For implementation details, see the [Amazon EventBridge API destinations documentation](#) and the [EventBridge connection authorization documentation](#).

#### QUESTION NO: 121

A company uses an Amazon DynamoDB table to store data that the company receives from devices. The DynamoDB table supports a customer-facing website to display recent activity on customer devices. The company configured the table with provisioned throughput for writes and reads.

The company wants to calculate performance metrics for customer device data on a daily basis. The solution must have minimal effect on the table's provisioned read and write capacity

Which solution will meet these requirements?

- A. Use an Amazon Athena SQL query with the Amazon Athena DynamoDB connector to calculate performance metrics on a recurring schedule.
- B. Use an AWS Glue job with the AWS Glue DynamoDB export connector to calculate performance metrics on a recurring schedule.
- C. Use an Amazon Redshift COPY command to calculate performance metrics on a recurring schedule.
- D. Use an Amazon EMR job with an Apache Hive external table to calculate performance metrics on a recurring schedule.

**ANSWER: B**

**Explanation:**

Use an AWS Glue job with the AWS Glue DynamoDB export connector to calculate performance metrics on a recurring schedule. The DynamoDB export connector uses DynamoDB's export-to-S3 capability to obtain table data for processing rather than reading items through normal DynamoDB data-plane operations. DynamoDB exports are performed from a point-in-time, consistent table snapshot and do not consume the table's provisioned read capacity or affect its write capacity. This isolates the daily analytics workload from the customer-facing application, whose reads and writes continue to use the table's configured provisioned throughput.

A scheduled AWS Glue job can process the exported data in Amazon S3 and calculate the required daily device metrics using scalable serverless Spark processing. DynamoDB point-in-time recovery must be enabled because exports rely on point-in-time recovery data. The company can schedule the Glue workflow or job daily, retain exported datasets as appropriate for reporting, and avoid introducing analytical scans against the production table. AWS documents that DynamoDB exports do not consume read capacity units and that the Glue DynamoDB export connector is intended to integrate DynamoDB exports into ETL workloads. See the [Amazon DynamoDB export to S3 documentation](#) and the [AWS Glue DynamoDB connection documentation](#).

**QUESTION NO: 122**

A company runs an application as a task in an Amazon ECS cluster. The application must have read and write access to a specific group of Amazon S3 buckets. The S3 buckets are in the same AWS Region and AWS account as the ECS cluster. The company needs to grant the application access to the S3 buckets according to the principle of least privilege.

Which combination of solutions will meet these requirements? (Select TWO.)

- A. Add a tag to each bucket. Create an IAM policy that includes a StringEquals condition that matches the tags and values of the buckets.
- B. Create an IAM policy that lists the full Amazon Resource Name (ARN) for each S3 bucket and the objects in each bucket.
- C. Attach the IAM policy to the instance role of the ECS task.
- D. Create an IAM policy that includes a wildcard Amazon Resource Name (ARN) that matches all combinations of the S3 bucket names.
- E. Attach the IAM policy to the task role of the ECS task.

**ANSWER: B E**

**Explanation:**

Creating an IAM policy that lists the full Amazon Resource Name (ARN) for each S3 bucket and its objects grants access only to the known buckets required by the application. For S3 permissions, the policy should scope bucket-level actions, such as `s3:ListBucket`, to each bucket ARN (`arn:aws:s3:::bucket-name`) and object-level read/write actions, such as `s3:GetObject` and `s3:PutObject`, to the corresponding object ARN (`arn:aws:s3:::bucket-name/*`). This

explicit resource scoping is a direct application of least privilege because the policy does not authorize access to unrelated buckets or objects.

Attaching that policy to the task role of the ECS task makes the permissions available to the application through the ECS task's AWS credentials. A task role is the IAM identity intended for application code running in containers, and ECS provides its temporary credentials to the task. This creates a workload-specific permission boundary, so the application receives only the S3 permissions defined for that task. AWS documents ECS task IAM roles at [Amazon ECS task IAM roles](#) and provides S3 action-to-resource policy examples in the [Amazon S3 identity-based policy examples](#).

### QUESTION NO: 123

The lead member of a DevOps team creates an AWS account. A DevOps engineer shares the account credentials with a solutions architect through a password manager application.

The solutions architect needs to secure the root user for the new account.

Which actions will meet this requirement? (Select TWO.)

- A. Update the root user password to a new, strong password.
- B. Secure the root user account by using a virtual multi-factor authentication (MFA) device.
- C. Create an IAM user for each member of the DevOps team. Assign the AdministratorAccess AWS managed policy to each IAM user.
- D. Create root user access keys. Save the keys as a new parameter in AWS Systems Manager Parameter Store.
- E. Update the IAM role for the root user to ensure the root user can use only approved services.

### ANSWER: A B

#### Explanation:

**Update the root user password to a new, strong password.** is required because the original root credentials have already been shared. Changing the password immediately revokes practical access based on the previously distributed password and establishes a password known only to the authorized root-user custodian. AWS recommends a unique, long, and strong password for the root user, which is reserved for a small number of account-level tasks that cannot be completed by other identities.

**Secure the root user account by using a virtual multi-factor authentication (MFA) device.** adds a separate verification factor to root-user sign-in. Consequently, possession of the password alone is insufficient to access the account as the root user. A virtual MFA device is an AWS-supported MFA type and is suitable for protecting the root user. AWS strongly recommends enabling MFA on every AWS account root user, particularly because root credentials have unrestricted access to account resources and billing information. After these measures, routine work should be performed using non-root identities under least-privilege permissions.

See AWS guidance on [root user best practices](#) and [enabling MFA for the root user](#).

### QUESTION NO: 124

A company needs to create an Amazon Elastic Kubernetes Service (Amazon EKS) cluster to host a digital media streaming application. The EKS cluster will use a managed node group that is backed by Amazon Elastic Block Store (Amazon EBS) volumes for storage. The company must encrypt all data at rest by using a customer managed key that is stored in AWS Key Management Service (AWS KMS)

Which combination of actions will meet this requirement with the LEAST operational overhead? (Select TWO.)

- A. Use a Kubernetes plugin that uses the customer managed key to perform data encryption.
- B. After creation of the EKS cluster, locate the EBS volumes. Enable encryption by using the customer managed key.

**C.** Enable EBS encryption by default in the AWS Region where the EKS cluster will be created. Select the customer managed key as the default key.

**D.** Create the EKS cluster. Create an IAM role with a policy that grants permission to use the customer managed key, and associate the role with the EKS cluster.

**E.** Store the customer managed key as a Kubernetes secret in the EKS cluster. Use the customer managed key to encrypt the EBS volumes.

**ANSWER: C D**

**Explanation:**

Enabling EBS encryption by default in the Region and selecting the customer managed KMS key ensures that EBS volumes created for the managed node group are automatically encrypted with that key. This applies encryption consistently at volume creation, including the root and attached EBS volumes that Amazon EKS creates for managed worker nodes, without requiring volume-by-volume actions or custom storage-encryption tooling. AWS documents that EBS encryption by default automatically encrypts newly created EBS volumes and snapshots using the configured default KMS key: [Amazon EBS encryption by default](#).

Creating the EKS cluster with an IAM role that is authorized to use the customer managed key allows Amazon EKS to use that key for envelope encryption of Kubernetes secrets stored in the cluster's control-plane data store. The appropriate KMS key permissions and key policy access enable EKS to create and manage the required KMS grant while keeping key management centralized in AWS KMS. Combining EKS secrets encryption with Regional EBS encryption by default covers the relevant control-plane secret data and managed-node EBS storage with minimal ongoing administration. Amazon EKS describes using a customer managed KMS key for Kubernetes secrets encryption in its [envelope encryption documentation](#).

**QUESTION NO: 125**

A company hosts an application on AWS Lambda functions that are invoked by an Amazon API Gateway API. The Lambda functions save customer data to an Amazon Aurora MySQL database. Whenever the company upgrades the database, the Lambda functions fail to establish database connections until the upgrade is complete. The result is that customer data is not recorded for some of the event.

A solutions architect needs to design a solution that stores customer data that is created during database upgrades.

Which solution will meet these requirements?

**A.** Provision an Amazon RDS proxy to sit between the Lambda functions and the database. Configure the Lambda functions to connect to the RDS proxy.

**B.** Increase the run time of the Lambda functions to the maximum. Create a retry mechanism in the code that stores the customer data in the database.

**C.** Persist the customer data to Lambda local storage. Configure new Lambda functions to scan the local storage to save the customer data to the database.

**D.** Store the customer data in an Amazon Simple Queue Service (Amazon SQS) FIFO queue. Create a new Lambda function that polls the queue and stores the customer data in the database.

**ANSWER: D**

**Explanation:**

Store the customer data in an Amazon Simple Queue Service (Amazon SQS) FIFO queue and use a Lambda function to poll the queue and write the data to the database. This design durably decouples API requests from database writes. When an API-invoked Lambda function receives customer data, it sends a message to Amazon SQS, which retains the message independently of Aurora availability. During a planned Aurora MySQL upgrade, the consumer Lambda function can be unable to connect to the database without losing the queued customer records. Once the upgrade is complete and connectivity is restored, Lambda continues polling the queue and processes the backlog.

Amazon SQS and Lambda provide at-least-once delivery semantics, so the database-writing function should be idempotent, for example by using a customer event ID or a unique database key to safely handle a possible repeated message delivery. A FIFO queue is appropriate when the application requires ordered processing within a message group and deduplication features. The queue's retention period should be configured to exceed the expected upgrade duration and recovery window. This architecture directly meets the requirement to preserve customer data created while the database is unavailable. See the [AWS Lambda documentation for using Amazon SQS as an event source](#) and the [Amazon SQS FIFO queue documentation](#).

#### QUESTION NO: 126

A company is building a mobile gaming app. The company wants to serve users from around the world with low latency. The company needs a scalable solution to host the application and to route user requests to the location that is nearest to each user.

Which solution will meet these requirements?

- A. Use an Application Load Balancer to route requests to Amazon EC2 instances that are deployed across multiple Availability Zones.
- B. Use a Regional Amazon API Gateway REST API to route requests to AWS Lambda functions.
- C. Use an edge-optimized Amazon API Gateway REST API to route requests to AWS Lambda functions.
- D. Use an Application Load Balancer to route requests to containers in an Amazon ECS cluster.

#### ANSWER: C

##### Explanation:

Use an edge-optimized Amazon API Gateway REST API to route requests to AWS Lambda functions. An edge-optimized API endpoint uses the Amazon CloudFront global edge network as the entry point for client requests. A user connects to a nearby CloudFront edge location, which reduces network latency for users distributed around the world before API Gateway processes and forwards the request to the API's configured AWS Region. This endpoint type is therefore well suited to a globally accessed mobile application when the API is hosted in a single Region. AWS Lambda provides serverless compute that automatically scales in response to incoming request volume, so it removes the need to provision, manage, or scale application servers. Together, API Gateway and Lambda provide a managed, scalable API hosting architecture while the edge-optimized endpoint delivers requests through a globally distributed network. API Gateway REST APIs support edge-optimized endpoint types, and AWS documents that these endpoints are designed for clients geographically distributed around the world. See the [API Gateway endpoint types documentation](#) and the [AWS Lambda scaling documentation](#).

#### QUESTION NO: 127

A company needs to allow a vendor to access CloudWatch Logs in the company's AWS account by using IAM roles for cross-account access.

Which solution will meet these requirements?

- A. Create roles in both accounts and trust the company role.
- B. Create a role in the vendor account and trust the company role.
- C. Create a role in the company account and trust the company role.
- D. Create a role in the company account with permissions and trust the vendor role.

#### ANSWER: D

##### Explanation:

Create a role in the company account with permissions and trust the vendor role. This follows the AWS cross-account IAM role pattern: the account that owns the protected resources creates the IAM role and attaches an identity-based permissions

policy granting only the required CloudWatch Logs actions, such as viewing log groups, streams, and events. The role's trust policy then identifies the vendor's AWS account or, preferably, the specific IAM role that the vendor will use as the trusted principal. A principal in the vendor account can call AWS STS AssumeRole to obtain temporary credentials for this role. Those temporary credentials operate in the company account and are limited by the permissions attached to the assumed role, allowing controlled access to the company's CloudWatch Logs resources without sharing long-term credentials. The vendor-side principal must also have permission to call `sts:AssumeRole` for the company role. Where a third party is involved, AWS recommends using an external ID condition in the trust policy to mitigate the confused-deputy risk. See the [AWS IAM cross-account role tutorial](#) and [AWS guidance on the confused deputy problem](#).

### QUESTION NO: 128

A company wants to send data from its on-premises systems to Amazon S3 buckets. The company created the S3 buckets in three different accounts. The company must send the data privately without traveling across the internet. The company has no existing dedicated connectivity to AWS.

Which combination of steps should a solutions architect take to meet these requirements? (Select TWO.)

- A.** Establish a networking account in the AWS Cloud. Create a private VPC in the networking account. Set up an AWS Direct Connect connection with a private VIF between the on-premises environment and the private VPC.
- B.** Establish a networking account in the AWS Cloud. Create a private VPC in the networking account. Set up an AWS Direct Connect connection with a public VIF between the on-premises environment and the private VPC.
- C.** Create an Amazon S3 interface endpoint in the networking account.
- D.** Create an Amazon S3 gateway endpoint in the networking account.
- E.** Establish a networking account in the AWS Cloud. Create a private VPC in the networking account. Peer VPCs from the accounts that host the S3 buckets with the VPC in the network account.

### ANSWER: A C

#### Explanation:

Establishing a networking account, creating a private VPC, and configuring AWS Direct Connect with a private virtual interface provides private IP connectivity from the on-premises environment to resources in that VPC. This establishes the required dedicated private network path into AWS rather than using an internet connection. An Amazon S3 interface endpoint in that VPC provides private connectivity to the S3 service through AWS PrivateLink. Interface endpoints use private IP addresses from the VPC subnets, so the on-premises network can reach the endpoint over the Direct Connect private virtual interface and send S3 requests without traversing the public internet.

The centralized networking VPC can provide this endpoint-based access for S3 buckets owned by multiple AWS accounts. Access to each bucket remains governed by IAM and S3 bucket policies; those policies can grant the required cross-account permissions and can restrict access to the specific VPC endpoint when appropriate. This architecture separates network connectivity from bucket ownership while retaining private routing end to end. AWS specifically documents that interface endpoints can be accessed from on-premises networks through Direct Connect or Site-to-Site VPN, subject to appropriate routing and DNS configuration. See [AWS PrivateLink concepts](#) and [Amazon S3 interface endpoints](#).

### QUESTION NO: 129

A company is using an Amazon Elastic Kubernetes Service (Amazon EKS) cluster. The company must ensure that Kubernetes service accounts in the EKS cluster have secure and granular access to specific AWS resources by using IAM roles for service accounts (IRSA).

Which combination of solutions will meet these requirements? (Select TWO.)

- A.** Create an IAM policy that defines the required permissions. Attach the policy directly to the IAM role of the EKS nodes.
- B.** Implement network policies within the EKS cluster to prevent Kubernetes service accounts from accessing specific AWS services.

C. Modify the EKS cluster's IAM role to include permissions for each Kubernetes service account. Ensure a one-to-one mapping between IAM roles and Kubernetes roles.

D. Define an IAM role that includes the necessary permissions. Annotate the Kubernetes service accounts with the Amazon Resource Name (ARN) of the IAM role.

E. Set up a trust relationship between the IAM roles for the service accounts and an OpenID Connect (OIDC) identity provider.

**ANSWER: D E**

**Explanation:**

IAM roles for service accounts (IRSA) provides AWS permissions at the Kubernetes service-account and pod level, rather than sharing the permissions of the Amazon EKS worker-node role. To use IRSA, an IAM role must contain the AWS permissions required by the workload and the Kubernetes service account must be annotated with that role's ARN. When a pod uses the annotated service account, the Amazon EKS pod identity webhook configures the pod to obtain temporary AWS credentials by assuming the specified role. This creates the required explicit association between the workload identity and narrowly scoped AWS permissions.

The IAM role also needs a trust relationship with the OpenID Connect (OIDC) identity provider associated with the EKS cluster. AWS Security Token Service validates the projected service-account token against this OIDC provider before issuing role credentials. For granular access, the trust policy should limit token claims such as `sub` to the intended Kubernetes namespace and service-account name, and typically limit `aud` to `sts.amazonaws.com`. Together, role annotation and OIDC-based trust allow each service account to assume only its intended IAM role and receive short-lived credentials. See the [Amazon EKS IRSA documentation](#) and AWS guidance for [associating an IAM role with a service account](#).

**QUESTION NO: 130 - (SIMULATION)**

A manufacturing company develops an application to give a small team of executives the ability to track sales performance globally. The application provides a real-time simulator in a popular programming language. The company uses AWS Lambda functions to support the simulator. The simulator is an algorithm that predicts sales performance based on specific variables.

Although the solution works well initially, the company notices that the time required to complete simulations is increasing exponentially. A solutions architect needs to improve the response time of the simulator.

Which solution will meet this requirement in the MOST cost-effective way?

**ANSWER: See the explanation for the answer**

**Explanation:**

**Answer: Configure AWS Lambda Provisioned Concurrency for the simulator Lambda function (on a published version or alias).**

Provisioned Concurrency keeps the configured number of Lambda execution environments initialized and ready to process simulator requests. This removes cold-start initialization delay and provides predictable, low-latency responses for the executive team's real-time simulator.

For a small, predictable user population, configure only the required number of provisioned concurrent executions (and optionally use scheduled scaling if usage follows business hours). This retains Lambda's operational simplicity and avoids the cost and management overhead of continuously running EC2, ECS/Fargate, or a batch-processing platform.

Publish a Lambda version.

Create or update an alias that points to that version.

Configure `Provisioned concurrency` on the alias/version with the number of warm execution environments needed.

AWS Batch is intended for asynchronous batch workloads and can add job-queue delay, so it is not appropriate for a real-time simulator.

### QUESTION NO: 131

A company plans to use an Amazon S3 bucket to archive backup data. Regulations require the company to retain the backup data for 7 years.

During the retention period, the company must prevent users, including administrators, from deleting the data. The company can delete the data after 7 years.

Which solution will meet these requirements?

- A.** Create an S3 bucket policy that denies delete operations for 7 years. Create an S3 Lifecycle policy to delete the data after 7 years.
- B.** Create an S3 Object Lock default retention policy that retains data for 7 years in governance mode. Create an S3 Lifecycle policy to delete the data after 7 years.
- C.** Create an S3 Object Lock default retention policy that retains data for 7 years in compliance mode. Create an S3 Lifecycle policy to delete the data after 7 years.
- D.** Create an S3 Batch Operations job to set a legal hold on each object for 7 years. Create an S3 Lifecycle policy to delete the data after 7 years.

### ANSWER: C

#### Explanation:

Create an S3 Object Lock default retention policy that retains data for 7 years in compliance mode. Create an S3 Lifecycle policy to delete the data after 7 years. This solution provides write-once-read-many protection for every new object uploaded to the Object Lock-enabled bucket. A default retention rule automatically applies a retain-until date seven years in the future, avoiding reliance on users or applications to set retention individually.

Compliance mode is specifically designed for regulatory retention requirements. Until an object's retention period expires, the object version cannot be overwritten or deleted by any principal, including users with administrative permissions and the AWS account root user. The retention date cannot be shortened after it has been applied. This supplies the required protection against deletion throughout the full seven-year period.

After the retain-until date passes, the Object Lock retention restriction no longer prevents expiration. An S3 Lifecycle expiration rule can then remove the eligible object versions automatically, allowing the company to delete archived data after its mandated retention period without manual intervention. The bucket must be created with S3 Object Lock enabled, which also enables versioning. See the [Amazon S3 Object Lock overview](#) and [Managing S3 Object Lock](#).

### QUESTION NO: 132

A company has implemented a self-managed DNS service on AWS. The solution consists of the following:

- Amazon EC2 instances in different AWS Regions
- Endpoints of a standard accelerator in AWS Global Accelerator

The company wants to protect the solution against DDoS attacks. What should a solutions architect do to meet this requirement?

- A.** Subscribe to AWS Shield Advanced. Add the accelerator as a resource to protect.
- B.** Subscribe to AWS Shield Advanced. Add the EC2 instances as resources to protect.
- C.** Create an AWS WAF web ACL that includes a rate-based rule. Associate the web ACL with the accelerator.
- D.** Create an AWS WAF web ACL that includes a rate-based rule. Associate the web ACL with the EC2 instances.

### ANSWER: A

#### Explanation:

Subscribe to AWS Shield Advanced and add the accelerator as a resource to protect. AWS Global Accelerator provides static anycast IP addresses and routes traffic across the AWS global network to healthy regional endpoints. AWS Shield Standard protections are automatically included with Global Accelerator, while AWS Shield Advanced provides enhanced DDoS detection and mitigation capabilities, access to the AWS Shield Response Team for qualifying events, and DDoS cost-protection benefits. A standard accelerator is a supported protected resource type for Shield Advanced, so protecting the accelerator applies DDoS protection at the internet-facing entry point that receives client DNS traffic before it is routed to the EC2 endpoints in different Regions. This architecture is especially appropriate for a multi-Region service because Global Accelerator centralizes the public ingress layer while continuing to direct traffic to healthy endpoints. Shield Advanced protection is configured by associating the Global Accelerator accelerator resource with the Shield Advanced subscription. For supported protected-resource types and the process for associating them, see the [AWS Shield Advanced Developer Guide](#). For details on Global Accelerator's routing and availability design, see the [AWS Global Accelerator Developer Guide](#).

#### QUESTION NO: 133

A solutions architect is designing a three-tier web application. The architecture consists of an internet-facing Application Load Balancer (ALB) and a web tier that is hosted on Amazon EC2 instances in private subnets. The application tier with the business logic runs on EC2 instances in private subnets. The database tier consists of Microsoft SQL Server that runs on EC2 instances in private subnets. Security is a high priority for the company. Which combination of security group configurations should the solutions architect use? (Select THREE.)

- A. Configure the security group for the web tier to allow inbound HTTPS traffic from the security group for the ALB.
- B. Configure the security group for the web tier to allow outbound HTTPS traffic to 0.0.0.0/0.
- C. Configure the security group for the database tier to allow inbound Microsoft SQL Server traffic from the security group for the application tier.
- D. Configure the security group for the database tier to allow outbound HTTPS traffic and Microsoft SQL Server traffic to the security group for the web tier.
- E. Configure the security group for the application tier to allow inbound HTTPS traffic from the security group for the web tier.
- F. Configure the security group for the application tier to allow outbound HTTPS traffic and Microsoft SQL Server traffic to the security group for the web tier.

#### ANSWER: A C E

##### Explanation:

Configure the security group for the web tier to allow inbound HTTPS traffic from the security group for the ALB. This ensures that web servers accept application requests only when they originate through the internet-facing load balancer, rather than directly from arbitrary network sources. AWS recommends restricting target security groups so that targets accept traffic only from their load balancer.

Configure the security group for the application tier to allow inbound HTTPS traffic from the security group for the web tier. Referencing the web-tier security group as the source provides a least-privilege, identity-based control for the tier-to-tier application connection, without relying on potentially changing instance IP addresses.

Configure the security group for the database tier to allow inbound Microsoft SQL Server traffic from the security group for the application tier. SQL Server commonly uses TCP port 1433, and limiting that database listener to the application-tier security group ensures that only the business-logic servers can initiate database connections. Security groups are stateful, so response traffic for permitted inbound connections is automatically allowed. These rules establish the intended ALB-to-web-to-application-to-database flow while minimizing exposure. See [AWS guidance for Application Load Balancer security groups](#) and [Amazon VPC security group documentation](#).

#### QUESTION NO: 134

A company's security team requests that network traffic be captured in VPC Flow Logs. The logs will be frequently accessed for 90 days and then accessed intermittently.

What should a solutions architect do to meet these requirements when configuring the logs?

- A. Use Amazon CloudWatch as the target. Set the CloudWatch log group with an expiration of 90 days
- B. Use Amazon Kinesis as the target. Configure the Kinesis stream to always retain the logs for 90 days.
- C. Use AWS CloudTrail as the target. Configure CloudTrail to save to an Amazon S3 bucket, and enable S3 Intelligent-Tiering.
- D. Use Amazon S3 as the target. Enable an S3 Lifecycle policy to transition the logs to S3 Standard-Infrequent Access (S3 Standard-IA) after 90 days.

**ANSWER: D**

**Explanation:**

Use Amazon S3 as the target. Enable an S3 Lifecycle policy to transition the logs to S3 Standard-Infrequent Access (S3 Standard-IA) after 90 days. VPC Flow Logs can publish flow-log records directly to an Amazon S3 bucket, providing durable, scalable storage for network-traffic records. Keeping the objects in S3 Standard for the first 90 days supports the stated frequent-access period, because S3 Standard is designed for frequently accessed data and provides low-latency retrieval.

After 90 days, an S3 Lifecycle configuration can automatically transition the objects to S3 Standard-IA. This storage class is intended for data accessed less frequently but that still requires millisecond retrieval when needed. It reduces storage cost while retaining the operational convenience of direct access, which fits the intermittent-access requirement. Lifecycle rules apply automatically based on object age, avoiding manual movement or administration of old log files. The S3 bucket can also serve as a long-term centralized repository for security analysis, audit retention, and integration with analytics services when required.

AWS documents S3 as a supported destination for VPC Flow Logs and describes the delivery configuration in the [Amazon VPC User Guide](#). S3 lifecycle transitions and the use case for S3 Standard-IA are documented in the [Amazon S3 User Guide](#).

**QUESTION NO: 135**

A company collects data from thousands of remote devices by using a RESTful web services application that runs on an Amazon EC2 instance. The EC2 instance receives the raw data, transforms the raw data, and stores all the data in an Amazon S3 bucket. The number of remote devices will increase into the millions soon. The company needs a highly scalable solution that minimizes operational overhead.

Which combination of steps should a solutions architect take to meet these requirements? (Select TWO.)

- A. Use AWS Glue to process the raw data in Amazon S3.
- B. Use Amazon Route 53 to route traffic to different EC2 instances.
- C. Add more EC2 instances to accommodate the increasing amount of incoming data.
- D. Send the raw data to Amazon Simple Queue Service (Amazon SQS). Use EC2 instances to process the data.
- E. Use Amazon API Gateway to send the raw data to an Amazon Kinesis data stream. Configure Amazon Kinesis Data Firehose to use the data stream as a source to deliver the data to Amazon S3.

**ANSWER: A E**

**Explanation:**

“Use Amazon API Gateway to send the raw data to an Amazon Kinesis data stream. Configure Amazon Kinesis Data Firehose to use the data stream as a source to deliver the data to Amazon S3.” provides a serverless, highly scalable ingestion path for millions of device requests. Amazon API Gateway can expose the REST API without managing web servers, while Amazon Kinesis Data Streams absorbs high-volume streaming records. Amazon Data Firehose can consume from a Kinesis data stream and reliably deliver the records to Amazon S3, eliminating the operational work of operating and scaling a fleet of ingestion and delivery instances. API Gateway supports direct AWS service integrations, including Kinesis actions, for an efficient managed integration path.

“Use AWS Glue to process the raw data in Amazon S3.” moves transformation into a managed, serverless data-integration service. AWS Glue can run scalable ETL jobs against data stored in Amazon S3, so compute resources are provisioned and managed by the service rather than by the company. This cleanly separates durable raw-data ingestion from downstream transformation, allowing each stage to scale independently and supporting batch or scheduled processing as data arrives. Together, these services provide managed ingestion, storage, and transformation with minimal operational overhead. See [Amazon API Gateway integration with Kinesis](#) and [What is AWS Glue](#).

#### QUESTION NO: 136

A gaming company has a web application that displays game scores. The application runs on Amazon EC2 instances behind an Application Load Balancer (ALB). The application stores data in an Amazon RDS for MySQL database.

Users are experiencing long delays and interruptions caused by degraded database read performance. The company wants to improve the user experience.

Which solution will meet this requirement?

- A. Use an Amazon ElastiCache (Redis OSS) cache in front of the database.
- B. Use Amazon RDS Proxy between the application and the database.
- C. Migrate the application from EC2 instances to AWS Lambda functions.
- D. Use an Amazon Aurora Global Database to create multiple read replicas across multiple AWS Regions.

#### ANSWER: A

##### Explanation:

Use an Amazon ElastiCache (Redis OSS) cache in front of the database. Game scores are typically read frequently, often by many users requesting the same or recently requested data. An ElastiCache for Redis OSS cluster provides an in-memory data store that the application can check before querying Amazon RDS for MySQL. When requested scores are present in the cache, the application returns them with very low latency and avoids a database read. This reduces read load on the RDS instance, helping restore responsive application performance during periods of high demand.

The application can use a cache-aside pattern: retrieve a score from Redis first, query RDS only on a cache miss, then populate or refresh the cache. It can also update or invalidate cached score data as new game results arrive. Redis OSS supports low-latency access and is well suited to caching application data that is repeatedly read. AWS specifically recommends ElastiCache to improve application and database performance by caching frequently accessed data. See the [Amazon ElastiCache User Guide](#) and [ElastiCache caching strategies documentation](#).

#### QUESTION NO: 137

A company is implementing a new policy to enhance the security of its AWS environment. The policy requires all administrative actions that users perform on the AWS Management Console to be secured by multi-factor authentication (MFA).

Which solution will allow the company to enforce this policy in the MOST operationally efficient way?

- A. Enable MFA on the root account. Ensure that all administrators use the root account to perform administrative actions.
- B. Create an IAM policy that requires MFA for administrative actions that administrators perform.
- C. Configure an Amazon CloudWatch alarm that sends an email notification when an administrator performs an administrative action without MFA.
- D. Use AWS Config to periodically audit IAM users and to automatically attach an IAM policy that requires MFA when AWS Config detects administrative actions.

#### ANSWER: B

### Explanation:

Create an IAM policy that requires MFA for administrative actions is the most operationally efficient solution because IAM evaluates authorization requests before allowing access to AWS resources and APIs. The policy can use the global condition key `aws:MultiFactorAuthPresent` to enforce MFA-based access. Commonly, administrators receive an explicit `Deny` for administrative actions when MFA is not present, while retaining only the limited permissions needed to enroll and manage their MFA device. This creates preventive enforcement rather than relying on a manual process or post-event detection.

The policy can be attached to the IAM identities or permission sets used by administrators, allowing the organization to consistently apply the requirement across its administrative population. When an administrator signs in to the AWS Management Console without completing MFA, requests for protected administrative operations are denied. After the administrator authenticates with MFA, the MFA context is present and the authorized administrative permissions can be used. This approach is centrally managed, scales with administrator onboarding and offboarding, and provides immediate enforcement at the access-control layer. AWS documents MFA enforcement with IAM policy conditions and provides example policies for denying requests made without MFA: [AWS global condition key: `aws:MultiFactorAuthPresent`](#) and [IAM policy examples for requiring MFA](#).

### QUESTION NO: 138

A company uses Amazon S3 to store customer data that contains personally identifiable information (PII) attributes. The company needs to make the customer information available to company resources through an AWS Glue Catalog. The company needs to have fine-grained access control for the data so that only specific IAM roles can access the PII data.

- A. Create one IAM policy that grants access to PII. Create a second IAM policy that grants access to non-PII data. Assign the PII policy to the specified IAM roles.
- B. Create one IAM role that grants access to PII. Create a second IAM role that grants access to non-PII data. Assign the PII policy to the specified IAM roles.
- C. Use AWS Lake Formation to provide the specified IAM roles access to the PII data.
- D. Use AWS Glue to create one view for PII data. Create a second view for non-PII data. Provide the specified IAM roles access to the PII view.

### ANSWER: C

### Explanation:

Use AWS Lake Formation to provide the specified IAM roles access to the PII data. Lake Formation is purpose-built to centrally govern data lakes that use Amazon S3 storage and the AWS Glue Data Catalog. It lets administrators grant permissions to IAM principals, including IAM roles, on catalog resources such as databases, tables, columns, and—in supported query engines—rows and cells. This makes it appropriate for protecting PII attributes while allowing authorized company resources to discover and query the same governed data through the Glue Catalog.

Lake Formation permissions are managed separately from broad S3 and Glue IAM permissions, enabling a consistent, centralized authorization model for data-lake access. For example, an administrator can grant a particular role access only to the columns that contain PII, or grant access to a table while excluding sensitive columns. Lake Formation integrates with AWS analytics services that use the Glue Data Catalog, allowing access controls to be enforced when approved principals access registered data-lake locations. This approach provides the required fine-grained controls without requiring separate copies of the customer dataset. See the [AWS Lake Formation Developer Guide](#) and the documentation on [data filters for row-level and cell-level security](#).

### QUESTION NO: 139

A company is migrating some of its applications to AWS. The company wants to migrate and modernize the applications quickly after it finalizes networking and security strategies. The company has set up an AWS Direct Connect connection in a central network account.

The company expects to have hundreds of AWS accounts and VPCs in the near future. The corporate network must be able to access the resources on AWS seamlessly and also must be able to communicate with all the VPCs. The company also wants to route its cloud resources to the internet through its on-premises data center.

Which combination of steps will meet these requirements? (Select THREE.)

- A.** Create a Direct Connect gateway in the central account. In each of the accounts, create an association proposal by using the Direct Connect gateway and the account ID for every virtual private gateway.
- B.** Create a Direct Connect gateway and a transit gateway in the central network account. Attach the transit gateway to the Direct Connect gateway by using a transit VIF.
- C.** Provision an internet gateway. Attach the internet gateway to subnets. Allow internet traffic through the gateway.
- D.** Share the transit gateway with other accounts. Attach VPCs to the transit gateway.
- E.** Provision VPC peering as necessary.
- F.** Provision only private subnets. Open the necessary route on the transit gateway and customer gateway to allow outbound internet traffic from AWS to flow through NAT services that run in the data center.

**ANSWER: B D F**

**Explanation:**

Creating a Direct Connect gateway and a transit gateway in the central network account, then using a transit virtual interface to connect the Direct Connect environment to the Direct Connect gateway, provides scalable connectivity between the corporate network and AWS. The Direct Connect gateway can be associated with the transit gateway, allowing routes for attached VPCs to be exchanged with the on-premises network over Direct Connect. This model avoids having to build and manage individual connectivity relationships for every VPC as the organization grows.

Sharing the transit gateway with other accounts and attaching VPCs to the transit gateway provides the hub-and-spoke connectivity model needed for hundreds of accounts and VPCs. AWS Resource Access Manager enables the central network account to share the transit gateway, while workload accounts create VPC attachments. The transit gateway then provides transitive routing among attached VPCs and to the Direct Connect-connected corporate network.

Provisioning only private subnets and adding the necessary Transit Gateway and customer gateway routes enables centralized internet egress through NAT services in the on-premises data center. Workload VPC route tables can direct default-route traffic to the transit gateway, which forwards it across Direct Connect to the on-premises network for NAT and internet access. See [AWS Direct Connect gateway associations for transit gateways](#) and [AWS Transit Gateway sharing](#).

**QUESTION NO: 140**

Question:

A genomics research company is designing a scalable architecture for a loosely coupled workload. Tasks in the workload are independent and can be processed in parallel. The architecture needs to minimize management overhead and provide automatic scaling based on demand.

Options:

- A.** Use a cluster of Amazon EC2 instances. Use AWS Systems Manager to manage the workload.
- B.** Implement a serverless architecture that uses AWS Lambda functions.
- C.** Use AWS ParallelCluster to deploy a dedicated high-performance cluster.
- D.** Implement vertical scaling for each workload task.

**ANSWER: B**

**Explanation:**

Implement a serverless architecture that uses AWS Lambda functions is the best fit because Lambda is designed to run independent, event-driven units of work without requiring the company to provision, patch, operate, or scale servers. Each task can be submitted as an invocation or event, and Lambda can run many executions concurrently, which aligns well with a loosely coupled workload whose tasks can be processed independently in parallel. AWS automatically manages the underlying compute capacity and scales the number of function environments in response to incoming demand, subject to the account's concurrency quotas and any configured reserved or provisioned concurrency settings. This removes most infrastructure-management activities while allowing the platform to scale down when no tasks are running, so the company pays primarily for invocations and execution duration. For genomics processing, the function implementation must remain within Lambda service limits, including the maximum execution duration, memory, ephemeral storage, and package or container-image constraints; individual tasks that fit these limits are a strong serverless use case. Lambda can also be integrated with event sources and services such as Amazon SQS to buffer and distribute asynchronous tasks reliably. See the [AWS Lambda Developer Guide](#) and [AWS Lambda concurrency documentation](#).

**QUESTION NO: 141**

An ecommerce company has noticed performance degradation of its Amazon RDS based web application. The performance degradation is attributed to an increase in the number of read-only SQL queries triggered by business analysts. A solutions architect needs to solve the problem with minimal changes to the existing web application.

What should the solutions architect recommend?

- A. Export the data to Amazon DynamoDB and have the business analysts run their queries.
- B. Load the data into Amazon ElastiCache and have the business analysts run their queries.
- C. Create a read replica of the primary database and have the business analysts run their queries.
- D. Copy the data into an Amazon Redshift cluster and have the business analysts run their queries

**ANSWER: C**

**Explanation:**

Create a read replica of the primary database and have the business analysts run their queries. Amazon RDS read replicas are asynchronously maintained, read-only copies of a source DB instance. They are specifically intended to scale read-heavy workloads by directing reporting, analytics, and other read-only query traffic away from the primary database. This reduces CPU, memory, storage I/O, and connection pressure on the primary instance, allowing it to focus on the application's transactional reads and writes.

This approach meets the requirement for minimal application changes because the existing primary RDS database and application write path remain in place. The analysts can instead connect to the read replica endpoint for their SQL queries. The replica uses the same database engine and contains a copy of the source data, so the analysts can generally continue using familiar SQL tooling and query patterns. Because replication is asynchronous, reports can reflect a small amount of replication lag; this is acceptable when the workload is read-only analysis rather than a requirement for immediate read-after-write consistency. AWS documents read replicas as a way to improve read performance and support reporting workloads. See [Working with read replicas for Amazon RDS](#) and [Amazon RDS read replica overview](#).

**QUESTION NO: 142**

A solutions architect is implementing a document review application using an Amazon S3 bucket for storage. The solution must prevent accidental deletion of the documents and ensure that all versions of the documents are available. Users must be able to download, modify, and upload documents.

Which combination of actions should be taken to meet these requirements? (Choose two.)

- A. Enable a read-only bucket ACL.
- B. Enable versioning on the bucket.
- C. Attach an IAM policy to the bucket.
- D. Enable MFA Delete on the bucket.

E. Encrypt the bucket using AWS KMS.

**ANSWER: B D**

**Explanation:**

Enabling versioning on the bucket preserves every version of an object when users upload a modified document using the same object key. Rather than overwriting the existing document irreversibly, Amazon S3 assigns a distinct version ID to the new upload and retains the prior version. This allows the application to maintain a complete version history and lets authorized users retrieve earlier document versions when needed. S3 versioning also changes the behavior of ordinary delete requests: deleting an object without specifying a version ID adds a delete marker, while the underlying prior versions remain stored and recoverable.

Enabling MFA Delete on the bucket adds protection against accidental or unauthorized permanent deletion of object versions. With MFA Delete enabled, permanently deleting a specified version or changing the bucket's versioning state requires the bucket owner's root credentials and a valid multi-factor authentication code. Users can still download documents and upload modified versions as part of normal document-review workflows, while the extra MFA requirement protects the retained historical versions from irreversible deletion. MFA Delete is available only for versioning-enabled buckets, so these controls work together to meet both retention and deletion-protection requirements. See [Amazon S3 Versioning](#) and [Configuring MFA Delete](#).

**QUESTION NO: 143**

A company has built an application that uses an Amazon Simple Queue Service (Amazon SQS) standard queue and an AWS Lambda function. The Lambda function writes messages to the SQS queue. The company needs a solution to ensure that the consumer of the SQS queue never receives duplicate messages.

Which solution will meet this requirement with the FEWEST changes to the current architecture?

- A. Modify the SQS queue to enable long polling for the queue.
- B. Delete the existing SQS queue. Recreate the queue as a FIFO queue. Enable content-based deduplication for the queue.
- C. Modify the SQS queue to enable content-based deduplication for the queue.
- D. Delete the SQS queue. Create an Amazon MQ message broker. Configure the broker to deduplicate messages.

**ANSWER: B**

**Explanation:**

Delete the existing SQS queue. Recreate the queue as a FIFO queue. Enable content-based deduplication for the queue. This is the smallest architectural change that provides Amazon SQS support for preventing duplicate messages from being delivered to the consumer. SQS FIFO queues provide exactly-once message processing: Amazon SQS does not introduce duplicate messages into the queue when a producer retries sending a message. When content-based deduplication is enabled, SQS calculates a SHA-256 hash of each message body and uses it as the deduplication ID. Messages with an identical body sent within the five-minute deduplication interval are accepted but are not delivered as separate messages. The Lambda producer can therefore continue sending messages to SQS, while the application changes its queue type from standard to FIFO. FIFO queues also require a message group ID on sent messages, which can be supplied by the Lambda function to preserve ordered processing within each group. For production resilience, consumers should still implement idempotent processing because any distributed consumer can retry work after a failure; however, FIFO queue deduplication is the applicable SQS mechanism for the stated requirement of avoiding duplicate queue messages. See the [Amazon SQS FIFO queues documentation](#) and [exactly-once processing documentation](#).

**QUESTION NO: 144**

A company wants to migrate an application that processes logs to AWS. Currently, the application runs on an on-premises storage area network (SAN).

The application reads and processes large log files sequentially. The application requires throughput of up to 500 MBps.

A solutions architect needs to migrate the application with minimal change to the application architecture.

Which solution will meet these requirements in the MOST cost-effective way?

- A. Use a General Purpose SSD (gp3) Amazon EBS volume with 500 MBps provisioned throughput.
- B. Use a Throughput Optimized HDD (st1) Amazon EBS volume with provisioned storage based on the throughput requirement.
- C. Use a Cold HDD (sc1) Amazon EBS volume with provisioned storage based on the throughput requirement.
- D. Use a Provisioned IOPS (io1) Amazon EBS volume with 500 MBps provisioned throughput.

**ANSWER: B**

**Explanation:**

Use a Throughput Optimized HDD (st1) Amazon EBS volume with provisioned storage based on the throughput requirement. Throughput Optimized HDD volumes are block-storage volumes designed specifically for frequently accessed, throughput-intensive workloads that perform large, sequential I/O, including log processing. This aligns directly with an application that reads large log files sequentially and needs sustained throughput rather than very high random IOPS or consistently low single-digit-millisecond latency.

An st1 volume can provide up to 500 MiB/s of throughput per volume, subject to the volume size and the attached EC2 instance's Amazon EBS bandwidth limits. Sizing the volume appropriately lets the workload reach the required throughput while using lower-cost HDD-backed EBS storage. EBS also maintains the block-device storage approach familiar to an application using an on-premises SAN, allowing it to be attached to an EC2 instance without requiring a fundamental application-storage redesign. AWS identifies st1 as suitable for big data, data warehouses, and log processing workloads that need high sequential throughput. See the [Amazon EBS volume types documentation](#) and [Amazon EBS pricing](#).

**QUESTION NO: 145**

A company creates a VPC that has one public subnet and one private subnet. The company attaches an internet gateway to the VPC. An Application Load Balancer (ALB) in the public subnet communicates with Amazon EC2 instances in the private subnet.

The EC2 instances in the private subnet must be able to download operating system and application updates from the internet. The instances must not be accessible from the internet.

Which combination of steps will meet these requirements? (Select THREE.)

- A. Associate an Elastic IP address with the NAT gateway.
- B. Add a route of 0.0.0.0/0 to the private subnet route table. Set the NAT gateway as a target.
- C. Deploy a NAT gateway in the public subnet.
- D. Deploy a NAT gateway in the private subnet.
- E. Add a route of 0.0.0.0/0 to the public subnet route table. Set the NAT gateway as a target.
- F. Associate an Elastic IP address with the internet gateway.

**ANSWER: A B C**

**Explanation:**

"Deploy a NAT gateway in the public subnet," "Associate an Elastic IP address with the NAT gateway," and "Add a route of 0.0.0.0/0 to the private subnet route table. Set the NAT gateway as a target" provide the standard AWS architecture for outbound-only internet access from private-subnet instances. A public NAT gateway is placed in a subnet whose route table sends internet-bound traffic to the VPC internet gateway. The Elastic IP address gives the NAT gateway a public IPv4 address that it uses when communicating with internet destinations. The EC2 instances themselves do not require public IPv4 addresses or direct routes to the internet gateway.

The private subnet's default route forwards traffic destined for external IPv4 addresses to the NAT gateway. The NAT gateway performs source network address translation, allowing response traffic for connections initiated by the instances to return through the gateway. This design permits the instances to retrieve operating system patches and application updates while preserving their private addressing and preventing unsolicited internet-initiated connections to them. The public Application Load Balancer can continue to send application traffic to the private instances through VPC-local networking and configured security groups.

AWS documents the required public-subnet placement and Elastic IP requirement for public NAT gateways in the [Amazon VPC NAT gateway basics](#) and describes routing private-subnet internet traffic through a NAT gateway in the [NAT gateway scenarios](#).

#### QUESTION NO: 146

A solutions architect needs to design a system to process incoming work items immediately. Processing can take up to 30 minutes and involves calling external APIs, executing multiple states, and storing intermediate states.

The solution must scale with variable workloads and minimize operational overhead.

Which combination of steps meets these requirements? (Select TWO.)

- A. Invoke an AWS Lambda function for each incoming work item. Configure each function to handle the work item completely. Store states in DynamoDB.
- B. Invoke an AWS Step Functions workflow to process incoming work items. Use Lambda functions for business logic. Store work item states in DynamoDB.
- C. Set up an API Gateway REST API to receive work items. Configure the API to invoke a Lambda function for each work item.
- D. Deploy two EC2 Reserved Instances behind an ALB and send requests to an SQS queue.
- E. Set up an API Gateway REST API to receive work items. Send the work items to an SQS queue.

#### ANSWER: B E

##### Explanation:

**Invoke an AWS Step Functions workflow to process incoming work items. Use Lambda functions for business logic. Store work item states in DynamoDB.** provides managed orchestration for the required multi-step processing. Step Functions coordinates individual business-logic tasks, supports state transitions, error handling, retries, and integrations with AWS services and external endpoints. A Standard workflow can run for up to one year, so it can manage a 30-minute end-to-end process while Lambda functions perform discrete, bounded units of work. DynamoDB can persist any application-specific intermediate data that must be retained or accessed independently of the workflow execution. Step Functions also retains workflow execution history and state-management capabilities, reducing the need to build and operate custom orchestration infrastructure. See [AWS Step Functions Developer Guide](#).

**Set up an API Gateway REST API to receive work items. Send the work items to an SQS queue.** creates a scalable, decoupled ingestion layer. API Gateway can accept incoming requests immediately, and Amazon SQS durably buffers work items while downstream processing scales according to demand. This separation lets the system absorb traffic bursts without requiring processing capacity to be available at the exact moment each request arrives. SQS integrates with AWS compute and event-driven services, enabling queued items to initiate the workflow-processing path while maintaining reliable asynchronous delivery. See [Amazon SQS Developer Guide](#).

#### QUESTION NO: 147

A software company needs to upgrade a critical web application. The application currently runs on a single Amazon EC2 instance that the company hosts in a public subnet. The EC2 instance runs a MySQL database. The application's DNS records are published in an Amazon Route 53 zone.

A solutions architect must reconfigure the application to be scalable and highly available. The solutions architect must also reduce MySQL read latency.

Which combination of solutions will meet these requirements? Select TWO.

- A. Launch a second EC2 instance in a second AWS Region. Use a Route 53 failover routing policy to redirect the traffic to the second EC2 instance.
- B. Create and configure an Auto Scaling group to launch private EC2 instances in multiple Availability Zones. Add the instances to a target group behind a new Application Load Balancer.
- C. Migrate the database to an Amazon Aurora MySQL cluster. Create the primary DB instance and reader DB instance in separate Availability Zones.
- D. Migrate the database to an Amazon RDS for MySQL Multi-AZ DB instance without a read replica.
- E. Place the current EC2 instance behind a Network Load Balancer and move the database to Amazon EBS gp2 volumes.

**ANSWER: B C**

**Explanation:**

Creating and configuring an Auto Scaling group to launch private EC2 instances in multiple Availability Zones, then adding the instances to a target group behind a new Application Load Balancer, provides a scalable and highly available web tier. The Application Load Balancer distributes requests across healthy targets, while the Auto Scaling group can maintain the desired capacity and replace unhealthy instances. Deploying the instances across multiple Availability Zones improves resilience to an Availability Zone failure. The Route 53 record can point users to the load balancer rather than to one individual EC2 instance. AWS recommends using Elastic Load Balancing with Amazon EC2 Auto Scaling for fault-tolerant, elastic applications. See [AWS documentation for Auto Scaling with load balancers](#).

Migrating the database to an Amazon Aurora MySQL cluster with a writer instance and a reader instance in separate Availability Zones addresses both database availability and read-latency requirements. Aurora replicas use the shared, distributed Aurora storage volume and can serve read-only traffic through the cluster reader endpoint. Applications can therefore direct read queries to reader instances, offloading the writer and reducing read-query latency under load. Aurora storage automatically maintains copies across multiple Availability Zones, and Aurora Replicas can support read scaling and failover. See [Amazon Aurora replication documentation](#).

**QUESTION NO: 148**

A company produces batch data that comes from different databases. The company also produces live stream data from network sensors and application APIs. The company needs to consolidate all the data into one place for business analytics. The company needs to process the incoming data and then stage the data in different Amazon S3 buckets. Teams will later run one-time queries and import the data into a business intelligence tool to show key performance indicators (KPIs).

Which combination of steps will meet these requirements with the LEAST operational overhead? (Choose two.)

- A. Use Amazon Athena for one-time queries Use Amazon QuickSight to create dashboards for KPIs
- B. Use Amazon Kinesis Data Analytics for one-time queries Use Amazon QuickSight to create dashboards for KPIs
- C. Create custom AWS Lambda functions to move the individual records from the databases to an Amazon Redshift cluster
- D. Use an AWS Glue extract transform, and load (ETL) job to convert the data into JSON format Load the data into multiple Amazon OpenSearch Service (Amazon Elasticsearch Service) clusters
- E. Use blueprints in AWS Lake Formation to identify the data that can be ingested into a data lake Use AWS Glue to crawl the source extract the data and load the data into Amazon S3 in Apache Parquet format

**ANSWER: A E**

**Explanation:**

Using blueprints in AWS Lake Formation to identify the data that can be ingested into a data lake, together with AWS Glue to crawl the source, extract the data, and load it into Amazon S3 in Apache Parquet format, provides a managed data-lake ingestion and transformation approach. Lake Formation blueprints can create the AWS Glue workflows that ingest data from sources such as databases and streaming services into an Amazon S3 data lake. AWS Glue then supplies serverless data

integration capabilities, including crawlers for cataloging data and jobs for transforming it. Apache Parquet is a columnar format that is well suited to analytics workloads because it can reduce data scanned and improve query efficiency in Amazon S3. See [AWS Lake Formation blueprints documentation](#) and [AWS Glue Parquet documentation](#).

Use Amazon Athena for one-time queries Use Amazon QuickSight to create dashboards for KPIs completes the analytics flow with minimal infrastructure management. Athena is a serverless interactive query service that uses SQL to analyze data directly in Amazon S3, making it appropriate for teams that need ad hoc, one-time analysis of the staged datasets. Amazon QuickSight can then connect to Athena query results or datasets to build and publish KPI visualizations and dashboards. This combination avoids managing query clusters while supporting the stated business-intelligence requirement.