

AWS CERTIFIED SOLUTIONS ARCHITECT - ASSOCIATE

SAA-C03

設計判断問題集

一問ずつ、条件を読み切る。

本番相当以上のシナリオ100問。

この問題集の使い方

昼休みや休憩時間に、1～3問ずつ解くための問題集です。サービス名を思い出すだけでは決まりません。問題文の制約を拾い、どの要件を優先するかを言葉にしてから選択してください。

答えを見る前に、棄却理由まで決める。

正解を一つ選んだあと、「残りは何の条件に反したか」を短く説明できれば、その問題は設計判断として身につきます。解答は必ず次ページから始まります。

1問の進め方

1. 数値要件、変更できないもの、最優先の評価軸に印を付ける。
2. 各選択肢が満たす条件と、落としている条件を一つずつ確かめる。
3. 最適解を決めてから次ページへ進み、判断過程まで照合する。

出題設計

公式試験ガイドの4ドメイン比率を100問へ置き換え、同じサービスでも制約によって結論が変わるように構成しています。

出題ドメイン	問題数
セキュアなアーキテクチャの設計	30問
レジリエントなアーキテクチャの設計	26問
高性能なアーキテクチャの設計	24問
コスト最適化アーキテクチャの設計	20問

難易度	問題数
本番相当	34問
本番より難しい	33問
複合難問	33問

本番相当 - 主要な制約を正しく拾えば、一つの設計へ収束します。

本番より難しい - 妥当な選択肢同士を、運用負荷やコストまで含めて比較します。

複合難問 - 複数分野を横断し、移行やDRなどの変更制約も判断に入れます。

QUESTION 001

問題1

SECURE

本番相当

単一選択

組織横断アクセスと緊急経路

ある金融会社は、AWS Organizationsの組織内に、本番、開発、監査用の80個のAWSアカウントを持っている。従業員は社内IdPで認証され、異動時の権限変更を1時間以内に全アカウントへ反映したい。各アカウントの管理者には業務上Administrator相当の権限が必要だが、CloudTrailの停止、許可されていないリージョンの利用、組織からの離脱は中央で禁止する。日常業務で管理アカウントを使うことは避けたい。一方、社内IdPまたはIAM Identity Centerを利用できない障害時にも、承認された担当者2名が重要なメンバーアカウントへ入り、復旧操作を実行できなければならない。長期アクセスキーは保管せず、緊急経路も定期的に試験して監査証跡を残す必要がある。これらの要件を最も適切に満たす構成はどれか。

- A** 組織インスタンスのIAM Identity Centerを社内IdPと連携し、グループへ権限セットを割り当てる。OUへSCPのガードレールを適用し、別系統の直接フェデレーションで限定IAM緊急ロールを引き受ける手順を用意して定期試験する。
- B** 各メンバーアカウントに同名の管理者IAMユーザーを作り、ハードウェアMFAと個別パスワードを配布する。組織管理アカウントにはCloudTrailを監視するロールだけを置き、権限制御は各アカウントのIAMポリシーへ委ねる。
- C** 各メンバーアカウントにIAM Identity Centerのアカウントインスタンスを作り、各管理者が権限セットを管理する。すべての制限を組織ルートSCPへ記述し、管理アカウントのユーザーと緊急ロールにも同じSCPが適用される前提で運用する。
- D** 組織インスタンスのIAM Identity Centerだけを唯一の認証経路とし、短いセッションとMFAを必須にする。障害時は既存セッションを持つ管理者が復旧し、セッションが全て失効した場合に限り管理アカウントのルートユーザーを共同利用する。

ここで答えを決める

正解・解説は次のページから

ANSWER 001

問題1の正解・解説

1 正解 A

- A** 組織インスタンスのIAM Identity Centerを社内IdPと連携し、グループへ権限セットを割り当てる。OUへSCPのガードレールを適用し、別系統の直接フェデレーションで限定IAM緊急ロールを引き受ける手順を用意して定期試験する。

2 問題文の重要な条件

- 80アカウントの権限を社内IdPから中央管理し、異動を短時間で反映する
- メンバーアカウントの管理者権限をSCPのガードレール内に制限する
- IdPまたはIAM Identity Centerの障害から独立した、監査可能な緊急アクセスが必要である
- 管理アカウントの日常利用と長期認証情報を避ける

3 正解へ至る判断過程

1. 多数のアカウントへの人のアクセスは、組織インスタンスのIAM Identity Centerと権限セットで一元化すると、グループ変更を横断的に反映できる。
2. Administrator権限を持つ利用者にも越えられない上限はSCPで設定する。ただしSCPは権限を付与せず、管理アカウントのユーザーやロールには作用しない点を分けて考える。
3. 通常のフェデレーションと同じ依存関係だけでは障害時に入れないため、限定アカウントのIAM緊急ロールへ別経路で直接フェデレーションする手順を事前構築する。
4. 緊急経路を複数人承認、短期セッション、ログ記録、定期試験で統制するAが、通常時と障害時の両方を満たす。

4 正解が要件を満たす理由

Aは、IAM Identity Centerの組織インスタンスで日常アクセスを一元管理し、SCPでメンバーアカウントの最大権限を制限する。さらに、通常系と障害原因を共有しない直接フェデレーション先の緊急IAMロールを限定的に用意するため、長期キーや管理アカウントの日常利用を増やさずに復旧経路を確保できる。緊急ロールを妨げないようSCPを設計し、利用を記録・試験することも要件に合う。

5・6 各選択肢の検討

不適切な理由と、正解になり得る条件を対にして確認します。

A 正解

中央のライフサイクル管理、メンバーアカウントのガードレール、通常系から独立した短期の緊急アクセスを同時に実現する。

B 不正解

MFAは認証を強化するが、80アカウントのIAMユーザーは異動時の迅速な無効化と棚卸しの運用負荷が大きく、各管理者が中央ガードレールを変更できる余地も残る。

この選択肢が正解になる条件

組織が数個の独立アカウントだけで、外部IdPを利用できず、中央ガードレールより各アカウントの完全な自治を優先する場合。

C 不正解

アカウントインスタンスは組織全体のAWSアカウントアクセスを一元管理する構成ではない。またSCPは管理アカウントのユーザーやロールには適用されないため、想定した統制にならない。

この選択肢が正解になる条件

単一アカウント内の対応アプリケーションだけに隔離ユーザーを与え、AWSアカウントへの組織横断アクセスを管理する必要がない場合。

D 不正解

IAM Identity CenterまたはIdPの停止時に新しいセッションを確立できず、既存セッション頼みとなる。管理アカウントのルートユーザー共同利用も最小権限と個人追跡性を損なう。

この選択肢が正解になる条件

緊急時にもIdPとIAM Identity Centerの可用性が契約上保証され、AWS側の独立した復旧経路を持たないリスクを組織が明示的に受容する場合。

7 今後使える判断基準

通常の人アクセスはIAM Identity Centerで集中管理し、権限付与は権限セット、越えられない組織上限はSCP、通常系障害時の入口は依存関係を分離した緊急フェデレーションとして別々に設計する。

8 関連サービスとの比較

IAM Identity Centerは人の割り当てと一時的なアカウントアクセスを集中管理する。SCPはメンバーアカウントの権限上限であって権限付与ではなく、管理アカウントには効かない。各アカウントのIAMユーザーは独立性が高い反面、認証情報と退職・異動処理が分散する。

QUESTION 002

問題2

RESILIENT

本番より難しい

単一選択

単一AZ障害点の解消

小売企業の注文APIは、2つのAZで有効化した1台のApplication Load Balancer（ALB）、EC2 Auto Scalingグループ、MySQLのAmazon RDS DBインスタンスで構成されている。ただし、ALBの全ターゲットを持つAuto ScalingグループのサブネットとRDSはAZ-aだけにあり、AZ-bには実行容量がない。EC2はステートレスで、通常3台、セール時は12台まで増える。先日のAZ-a障害では正常ターゲットとDBを同時に失い、API全体が2時間停止した。会社は同一リージョン内の単一AZ障害に耐え、月間99.95%を目標とし、DBのコミット済み更新を失わず、5分以内に復旧したい。アプリケーションやDBエンジンの変更は今期行えず、平常時の読み取り性能増強も不要であり、既存の公開DNS名も維持したい。運用作業と追加コストを要件達成に必要な範囲へ抑える構成はどれか。

- A** 2つの独立スタックを作り、各スタックのALB自体は2AZで有効化するが、プライマリのEC2はAZ-a、セカンダリのEC2はAZ-bに固定する。Route 53でALBを切り替え、DBは別AZの非同期RDSリードレプリカを手動昇格する。
- B** 1台のALBで少なくとも2つのAZを有効化し、Auto Scalingグループを両AZのプライベートサブネットへ分散してELBヘルスチェックを使う。RDSを同期スタンバイを持つMulti-AZ DBインスタンス配置へ変更する。
- C** Auto Scalingグループだけを2つのAZへ拡張し、既存ALBはそのまま利用するがRDSはAZ-aに残す。EC2の最小容量を各AZ3台相当に増やし、RDSの自動スナップショット間隔を短くして障害時に別AZへ復元する。
- D** 別リージョンにALB、Auto Scalingグループ、RDSクロスリージョンリードレプリカを常時配置する。Route 53でリージョンを切り替え、障害時にレプリカを昇格してからアプリケーションの接続先を変更する。

ここで答えを決める

正解・解説は次のページから

ANSWER 002

問題2の正解・解説

1 正解 **B**

- B** 1台のALBで少なくとも2つのAZを有効化し、Auto Scalingグループを両AZのプライベートサブネットへ分散してELBヘルスチェックを使う。RDSを同期スタンバイを持つMulti-AZ DBインスタンス配置へ変更する。

2 問題文の重要な条件

- 同一リージョンの単一AZ障害をアプリケーション層とDB層の双方で吸収する
- コミット済みDB更新を失わず、RTOを5分以内にする
- アプリケーションとDBエンジンは変更できない
- 読み取り拡張やリージョン障害対策は要求されず、コストを必要範囲に抑える

3 正解へ至る判断過程

1. リクエスト経路、EC2容量、DBのいずれかが1つのAZだけに残れば、AZ障害点は解消されない。3層すべてを確認する。
2. ALBとAuto Scalingグループを複数AZに構成すれば、正常なAZのEC2ヘルレーティングし、不足インスタンスを自動補充できる。
3. RDS Multi-AZ DBインスタンス配置は別AZの同期スタンバイへ自動フェイルオーバーするため、非同期リードレプリカの手動昇格よりRPOとRTOに合う。
4. マルチリージョンは範囲外の障害にも備えられるが、コストと運用が過剰なのでBを選ぶ。

4 正解が要件を満たす理由

Bでは、ALB、Auto Scalingグループ、RDSの各層から単一AZ依存を除去する。ALBは正常なターゲットだけへ転送し、Auto Scalingは両AZで必要容量を維持する。RDS Multi-AZは別AZへ同期的に複製したスタンバイへ自動フェイルオーバーするため、アプリケーションのDBエンジン変更や読み取り分散を伴わず、短いRTOとコミット済み更新を失わない要件に最も適する。

5・6 各選択肢の検討

不適切な理由と、正解になり得る条件を対にして確認します。

A 不正解

Web層はAZ分離できるが、通常のリードレプリカは非同期で、昇格も手動である。障害直前の更新損失と5分RTOを確実に満たせない。

この選択肢が正解になる条件

多少のレプリカ遅延によるデータ損失と手動切り替えを許容し、待機系をレポート読み取りにも利用したい場合。

B 正解

受付、コンピュート、データベースをすべて複数AZ化し、同期スタンバイと自動フェイルオーバーを最小限の構成変更で利用できる。

C 不正解

EC2は分散されてもRDSがAZ-aに残るため、DB障害でAPIは停止する。スナップショット復元はRPOが取得時点まで戻り、5分以内の復旧にも向かない。

この選択肢が正解になる条件

EC2ホスト障害だけを対象とし、ALBとDBの停止は許容でき、RTOが数時間、RPOが最新スナップショット時点でよい場合。

D 不正解

リージョン障害まで保護できるが、非同期レプリカの昇格と接続変更が必要で、提示されたRPOを保証しにくい。単一AZ対策として待機リージョンの費用と運用も過剰である。

この選択肢が正解になる条件

リージョン全体の災害が対象で、定義済みの非ゼロRPOとDNS切り替えを含む長めのRTOを受容し、追加コストを負担できる場合。

7 今後使える判断基準

AZ耐障害性は、ロードバランサー、実行容量、状態保存層を端から順にたどり、1層でも単一AZに残っていないかを確認する。RPOほぼゼロと自動切り替えには、読み取り用の非同期レプリカではなく同期HA配置を選ぶ。

8 関連サービスとの比較

RDS Multi-AZ DBインスタンスは高可用性の同期スタンバイで、通常は読み取り先ではない。RDSリードレプリカは読み取り拡張と非同期DRに向く。ALBを複数AZで有効にしても、正常ターゲットやDBが単一AZだけならアプリケーション全体はAZ障害に耐えない。

QUESTION 003

問題3

PERFORMANCE

複合難問

単一選択

Aurora読み取り接続の弾力化

動画配信会社はAurora PostgreSQLクラスターを使用し、1台のライターと2台の同一クラスのAuroraレプリカを別々のAZに配置している。APIは書き込みをクラスターエンドポイントへ送る一方、読み取り接続文字列には1台目のレプリカのインスタンスエンドポイントを設定している。新作公開時には読み取り接続が通常の8倍となり、1台目のCPUが95%に達してp95応答時間が2秒を超えるが、2台目はほぼアイドルである。この間もライターのCPUと書き込み遅延、クラスターのストレージI/Oには十分な余裕がある。負荷は30～90分で収まり、平常時の余剰コストは避けたい。読み取りレプリカを最低2台、最大8台の範囲で自動増減し、追加されたレプリカを将来のアプリケーション配布なしに利用させたい。既存のSQLと書き込み経路は変更しない。最も適切な改善策はどれか。

- A** 読み取り接続を既存2台の各インスタンスエンドポイントへ半分ずつ静的に割り当て、両方を最大想定負荷に耐えるクラスへスケールアップする。負荷終了後もクラスを維持し、障害時だけ設定を手動変更する。
- B** 読み取り用のNLBを作成し、2台のAuroraレプリカのIPアドレスをターゲットとして登録する。NLBのヘルスチェックで接続を分散し、CloudWatchアラーム時に運用担当者がレプリカとターゲットを追加する。
- C** 読み取り接続をAuroraのreader endpointへ変更し、平均CPU使用率または平均接続数を指標とするAurora Auto Scalingを最小2、最大8レプリカで設定する。接続プールはDNS変更後に再接続できるよう調整する。
- D** 2台のレプリカを削除してライターをAurora Serverless v2へ変更し、読み書きの全接続をクラスターエンドポイントへ統合する。最大ACUをピークに合わせ、読み取りの増加はライターの垂直スケーリングだけで処理する。

ここで答えを決める

正解・解説は次のページから

ANSWER 003

問題3の正解・解説

1 正解 C

- C** 読み取り接続をAuroraのreader endpointへ変更し、平均CPU使用率または平均接続数を指標とするAurora Auto Scalingを最小2、最大8レプリカで設定する。接続プールはDNS変更後に再接続できるよう調整する。

2 問題文の重要な条件

- 現在は1台のインスタンスエンドポイントへ読み取りが集中している
- 短時間の8倍負荷に合わせてレプリカ数を2～8台で自動調整する
- 自動追加されたレプリカをアプリケーション再配布なしで利用する
- SQLとライターへの接続経路は変更せず、平常時コストを抑える

3 正解へ至る判断過程

1. まず容量不足だけでなく接続先の偏りを特定する。インスタンスエンドポイントは指定した1台に固定されるため、既存の2台目すら活用できていない。
2. reader endpointは利用可能なAuroraレプリカ間で新しい接続を分散し、Auto Scalingで追加されたレプリカも接続先集合へ含められる。
3. Aurora Auto Scalingの目標追跡で最小・最大レプリカ数を指定すれば、一時的な読み取り増加後にスケールインして余剰費用を抑えられる。
4. reader endpointはクエリ単位ではなく接続単位で分散するため、長寿命接続を更新できる接続プール設定まで含むCが最適である。

4 正解が要件を満たす理由

Cは、読み取り接続の偏りと容量の弾力性という二つの原因を同時に解決する。reader endpointは利用可能なレプリカへ接続を分散し、Aurora Auto Scalingは指定した指標と2～8台の範囲でレプリカを増減する。アプリケーションは固定のreader endpointを使い続けられるため、新規レプリカごとの設定配布が不要であり、書き込み経路も変わらない。再接続を促す設定は接続単位の分散という性質にも対応する。

5・6 各選択肢の検討

不適切な理由と、正解になり得る条件を対にして確認します。

A 不正解

既存2台への偏りは減るが、自動追加したレプリカを発見できず、ピーク用の大きなクラスを常時維持するため、弾力性と平常時コストの要件を満たさない。

この選択肢が正解になる条件

負荷が常時ほぼ一定で2台から増減せず、アプリケーション側で接続先ごとの重みと障害切り替えを確実に管理できる場合。

B 不正解

Auroraの管理対象エンドポイントを使わず、DBインスタンスの変化をNLBターゲットへ追従させる独自運用が必要になる。手動追加はピークへの自動対応要件にも反する。

この選択肢が正解になる条件

Aurora以外の固定IP対応TCPサービスを分散し、ターゲットの増減がまれで、独自の登録自動化を運用する必要がある場合。

C 正解

reader endpointによる接続分散とAurora Auto Scalingによるレプリカ数の弾力化を組み合わせ、固定接続先のままピークと平常時コストに対応する。

D 不正解

ライターだけのスケールでは読み取りを独立して水平分散できず、書き込みと読み取りが同じ計算資源を競合する。既存SQLを変えなくても可用性と読み取り弾力性が後退する。

この選択肢が正解になる条件

総負荷が小さく断続的で、単一のServerless v2ライターの範囲内に収まり、読み取り分離や即時フェイルオーバー用レプリカが不要な場合。

7 今後使える判断基準

Auroraの読み取り拡張では、レプリカを増やす前にアプリケーションがreader endpointを使っているかを確認する。Auto Scalingは台数を変える仕組み、reader endpointは変化する接続先を隠蔽する仕組みであり、両方が必要である。

8 関連サービスとの比較

instance endpointは特定DBインスタンスへの診断や専用ワークロードに向く。reader endpointはレプリカ群へ接続を分散するが、個々のクエリを再分配はしない。Aurora Auto Scalingはレプリカ数を変え、Serverless v2は各DBインスタンスの容量を細かく変える。

QUESTION 004

問題4

COST

本番相当

複数選択 (2つ)

変更予定を織り込む購入割引

ある企業は過去12か月のコンピュータ利用を分析した。全体の約55%は24時間稼働するLinuxのEC2とFargateで、今後18か月にEC2のインスタンスファミリー、サイズ、リージョンを段階的に変更し、一部をFargateとLambdaへ移す予定である。約15%はライセンス認証装置を動かす12台のEC2で、今後3年間、同じインスタンスタイプ、OS、テナンシー、特定AZで常時稼働する。この装置は再起動後も同AZですぐ起動できる容量確保が必要である。残り30%は季節変動が大きく、確約対象にしたくない。現在のオンデマンド利用は組織の一括請求に集約され、購入後も月次で使用率とカバレッジを監視できる。財務部は過去の最低利用量を超えて契約せず、割引を最大化しながら移行を妨げない購入方針を求めている。採用すべき方針を2つ選択してください。

- A** 変化する55%のうち、履歴上確実な時間当たり利用額だけをCompute Savings Plansで1年確約する。EC2のファミリー、サイズ、OS、テナンシー、リージョン変更やFargate、Lambdaへの移行にも割引を追従させる。
- B** 固定された15%には、対象AZ、インスタンスタイプ、プラットフォーム、テナンシーに一致するゾーンのStandard Reserved Instancesを購入し、割引に加えてそのAZのキャパシティ予約を得る。
- C** 55%と15%をまとめてEC2 Instance Savings Plansで3年確約する。現在最も多いリージョンとインスタンスファミリーを指定すれば、FargateとLambda、および別リージョンへ移行した利用にも同じ割引が自動適用される。
- D** 季節ピークを含む利用全体の70%に対し、現在のEC2構成と一致するリージョンスキューのStandard Reserved Instancesを3年分購入する。余剰分は将来別ファミリー、別OS、別リージョンへ自由に交換する。

ここで答えを決める

正解・解説は次のページから

ANSWER 004

問題4の正解・解説

1 正解 **A・B**

- A** 変化する55%のうち、履歴上確実な時間当たり利用額だけをCompute Savings Plansで1年確約する。EC2のファミリー、サイズ、OS、テナンシー、リージョン変更やFargate、Lambdaへの移行にも割引を追従させる。
- B** 固定された15%には、対象AZ、インスタンスタイプ、プラットフォーム、テナンシーに一致するゾーンスコープのStandard Reserved Instancesを購入し、割引に加えてそのAZのキャパシティ予約を得る。

2 問題文の重要な条件

- 55%の基礎負荷はEC2構成とリージョンが変わり、一部はFargateとLambdaへ移る
- 15%は3年間構成とAZが固定され、割引だけでなく起動容量の確保が必要である
- 変動する30%と不確実な利用額を確約に含めない
- 異なる性質の負荷へ一つの購入モデルを無理に当てはめない

3 正解へ至る判断過程

1. 確約割引は、予測可能な最低利用量と、今後変わらない属性の範囲を分けて評価する。
2. Compute Savings Plansは時間当たり利用額の確約で、EC2の構成・リージョン変更だけでなくFargateとLambdaへの移行にも柔軟なので、変化する基礎負荷に合う。
3. Savings Plans自体はキャパシティを予約しない。特定AZで構成が固定され容量保証が必要な装置には、条件が一致するゾーンRIが適する。
4. 季節ピークや移行後に消える可能性のある額まで長期確約しないAとBの組み合わせが最適である。

4 正解が要件を満たす理由

Aは、移行予定のある基礎負荷について、確実に消費する金額だけをCompute Savings Plansへ載せるため、EC2の構成やリージョン変更、Fargate・Lambda化を妨げない。Bは、3年間変わらない装置に限ってゾーンStandard RIを使い、強い割引と特定AZのキャパシティ予約を同時に得る。柔軟性が必要な部分と容量保証が必要な固定部分を分け、変動30%をオンデマンドのまま残す点が全条件に合う。

5・6 各選択肢の検討

不適切な理由と、正解になり得る条件を対にして確認します。

A 正解

Compute Savings Plansの広い適用柔軟性を、構成変更とFargate・Lambda移行がある予測可能な最低利用額に限定して使っている。

B 正解

構成もAZも固定された負荷に限るためStandard RIの制約と一致し、ゾーンスコープで必要なキャパシティ予約も得られる。

C 不正解

EC2 Instance Savings Plansは指定リージョンの特定インスタンスファミリーに結び付く。Fargate、Lambda、別リージョンへの移行分まで同じ割引が追隨するという前提が誤りで、3年確約も変更計画に対して硬すぎる。

この選択肢が正解になる条件

利用が指定リージョン内の同一EC2ファミリーに3年間安定し、Fargate、Lambda、別リージョンへ移行せず、より高い割引率を優先する場合。

D 不正解

Standard RIは別インスタンスファミリー、別OS、別リージョンへ自由に交換できない。季節ピークまで70%を3年確約すると、未使用割引のリスクも高い。

この選択肢が正解になる条件

対象負荷が同一リージョン、同一ファミリーとプラットフォームで3年間確実に継続し、確約量が季節を通じた最低稼働以下である場合。

7 今後使える判断基準

確約割引は、まず確実な最低利用量だけを対象にし、構成変更の幅でCompute Savings Plans、EC2ファミリー固定ならEC2 Instance Savings Plans、特定AZの容量確保まで必要ならゾーンRIを選び分ける。

8 関連サービスとの比較

Compute Savings PlansはEC2のファミリー・リージョンをまたぎFargateとLambdaにも柔軟だが、容量を予約しない。EC2 Instance Savings Plansはリージョンとファミリーが固定される代わりに割引が強い。ゾーンRIは属性一致が必要だがキャパシティ予約を含む。

QUESTION 005

問題5

SECURE

本番より難しい

単一選択

インターネット非接続のSSM運用

医療企業は、個人情報処理用のEC2をプライベートサブネットに配置し、インターネットゲートウェイ、NAT Gateway、踏み台、インバウンドSSHを一切使わない方針である。最新のSSM Agentを組み込んだ標準AMIとSystems Manager用のインスタンスプロファイルは用意済みである。運用担当者はSession ManagerとRun Commandを使用し、Patch ManagerでOSを更新する。セッションとコマンドのログはCloudWatch Logsへ、出力ファイルは社内S3バケットへ保存し、Session ManagerにはカスターマネージドKMSキーの暗号化を使用する。Amazon提供のパッチ関連S3コンテンツにも到達させたい。通信はAWSネットワーク内に限定し、必要な経路だけを最小権限で許可する。最も適切なネットワーク構成はどれか。

- A** Systems Manager用のGateway VPC endpointを1つ作成し、プライベートサブネットのルートテーブルへ登録する。S3、CloudWatch Logs、KMSへの通信は同じエンドポイントが中継するため、追加エンドポイントは作成しない。
- B** ssmとssmmessagesのInterface VPC endpointを作成し、Private DNSを有効にする。インスタンスのセキュリティグループで443アウトバウンドだけを許可すれば、ログ、KMS暗号化、パッチ取得も自動的に同じ経路を通る。
- C** 管理専用のパブリックサブネットにNAT Gatewayを置き、プライベートサブネットから0.0.0.0/0をNATへ向ける。AWSサービスの公開エンドポイントだけをDNS名で許可し、Session Manager利用時はSSHポートを閉じる。
- D** ssmとssmmessagesのInterface VPC endpoint、S3のGateway endpoint、CloudWatch LogsとKMSのInterface endpointを作りPrivate DNSを有効にする。endpoint側SGはEC2からの443を許可し、S3ポリシーとIAMを必要なバケットへ限定する。

ここで答えを決める

正解・解説は次のページから

ANSWER 005

問題5の正解・解説

1 正解 **D**

- D** ssmとssmmessagesのInterface VPC endpoint、S3のGateway endpoint、CloudWatch LogsとKMSのInterface endpointを作りPrivate DNSを有効にする。endpoint側SGはEC2からの443を許可し、S3ポリシーとIAMを必要なバケットへ限定する。

2 問題文の重要な条件

- EC2にはNAT、インターネット経路、踏み台、インバウンド管理ポートを持たせない
- Session Manager、Run Command、Patch Managerを最新SSM Agentで利用する
- CloudWatch Logs、社内およびAWS管理S3バケット、KMSへプライベートに到達する
- エンドポイントのセキュリティグループ、IAM、ポリシーを最小権限にする

3 正解へ至る判断過程

1. SSM Agentは外向き接続を開始するためEC2へのインバウンド管理ポートは不要だが、NATなしでは各AWSサービスへのプライベート経路が必要になる。
2. Systems Managerの制御とデータチャネルにはssmおよびssmmessagesのInterface endpointを使い、endpoint側のセキュリティグループでEC2からのHTTPSを受ける。
3. ログ、KMS、S3は別サービスなので、それぞれLogsとKMSのInterface endpoint、S3 Gateway endpointを用意する。
4. Private DNS、S3のAWS管理バケットへの必要権限、インスタンスプロファイルまで含むDだけが、全機能をインターネットなしで成立させる。

4 正解が要件を満たす理由

Dは、SSM Agentの制御通信とSession Managerのデータチャネルをssm・ssmmessagesのInterface endpointへ閉じ、ログとKMSにも個別のInterface endpoint、パッチ資材と出力用S3にはGateway endpointを提供する。Interface endpointのセキュリティグループはEC2側からの443を受ける必要があり、Private DNSで通常のサービス名をプライベートIPへ解決する。IAMとS3 endpoint policyを絞ることで、到達性と最小権限を両立する。

5・6 各選択肢の検討

不適切な理由と、正解になり得る条件を対にして確認します。

A 不正解

Systems ManagerはGateway endpoint型ではなくInterface endpointを使用する。また1つのエンドポイントがS3、Logs、KMSを代理することではなく、各サービスへの経路が欠ける。

この選択肢が正解になる条件

対象がGateway endpointを提供するS3またはDynamoDBだけで、Systems ManagerやCloudWatch Logs、KMSを利用しない場合。

B 不正解

Session ManagerとRun Commandの主要通信は可能になるが、NATがないためLogs、KMS、S3への経路が不足する。Interface endpoint側SGの443インバウンド許可も明示されていない。

この選択肢が正解になる条件

ログ出力、KMS暗号化、S3上のファイルやパッチ資材を使わず、ssmとssmmessagesだけで完結する操作に限定する場合。

C 不正解

機能上は公開エンドポイントへ到達できるが、インターネット経路とNATを一切使わない明示要件に反し、宛先をDNS名だけで厳格に制御する運用も複雑である。

この選択肢が正解になる条件

アウトバウンドのインターネットアクセスが許可され、複数Interface endpointの固定費よりNAT経由の共通出口を優先する場合。

D 正解

利用するAWSサービスごとに必要なPrivateLinkまたはGateway endpointを揃え、DNS、SG、IAM、S3ポリシーまで含めて閉域運用を成立させる。

7 今後使える判断基準

NATなしのSystems Manager問題では、SSM本体だけでなく実際に使う機能の依存先を列挙する。ssm・ssmmessages、S3、Logs、KMSなどをサービスごとにendpointへ対応させ、endpoint側SGの受信443とPrivate DNSも確認する。

8 関連サービスとの比較

Interface endpointはENIとセキュリティグループを持ち、Systems Manager、Logs、KMSなどをPrivateLinkで提供する。S3 Gateway endpointはルートテーブルとendpoint policyで制御する。NATは多くの公開サービスへ出られるが、インターネット非接続要件には合わない。

QUESTION 006

問題6

RESILIENT

複合難問

単一選択

長時間SQS処理の再実行制御

画像処理会社はAmazon SQS StandardキューからEC2ワーカーがジョブを取得し、処理結果と課金レコードをデータベースへ書き込んでからメッセージを削除している。処理時間の中央値は45秒だが、大きな画像では12分かかる。visibility timeoutは60秒のため、処理中のメッセージが別ワーカーへ再表示され、同じ注文が二重課金されることがある。また、破損画像は何度再試行しても失敗し、正常ジョブを圧迫している。ピーク時は最大1,000台のワーカーが並列稼働するため、重複処理は計算費用も大きく増やす。1日500万件の順不同処理を維持し、一時的なワーカー障害ではジョブを失わず再実行し、破損ジョブは調査後に再投入したい。データベースは注文IDに対する条件付き書き込みを利用できる。運用負荷を抑えつつ最も堅牢にする変更はどれか。

- A** visibility timeoutを通常の最大処理時間に合わせ、長時間ジョブはChangeMessageVisibilityで延長する。成功コミット後だけ削除し、注文IDの条件付き書き込みで冪等化する。maxReceiveCount付きDLQとredriveを設定する。
- B** キューをSQS FIFOへ変更し、content-based deduplicationだけで二重課金を防ぐ。visibility timeoutは60秒のままにし、破損画像は同じMessageGroupIdの末尾で成功するまで無期限に再試行する。
- C** visibility timeoutを0秒にして全ワーカーが同じジョブを並列実行し、最初に完了したワーカーだけがメッセージを削除する。データベースの重複レコードは毎晩のバッチで検出して返金する。
- D** メインキューの保持期間を最大にし、maxReceiveCountを1として全失敗を直ちにDLQへ移す。二重処理を避けるため、ワーカーはメッセージ受信直後、DB更新より前にキューから削除する。

ここで答えを決める

正解・解説は次のページから

ANSWER 006

問題6の正解・解説

1 正解 **A**

- A** visibility timeoutを通常の最大処理時間に合わせ、長時間ジョブはChangeMessageVisibilityで延長する。成功コミット後だけ削除し、注文IDの条件付き書き込みで冪等化する。
maxReceiveCount付きDLQとredriveを設定する。

2 問題文の重要な条件

- 処理時間が60秒を大きく超え、最大12分まで変動する
- Standardキューの少なくとも1回配送とワーカー障害時の再実行を受け入れる
- 課金副作用は重複させず、破損ジョブは隔離して後から再投入する
- 高スループットの順不同処理を維持し、注文IDの条件付き書き込みを使える

3 正解へ至る判断過程

1. 再表示の直接原因はvisibility timeoutが処理時間より短いことである。十分な初期値と処理中の延長で、正常処理中の不要な並行実行を減らす。
2. それでもStandardキューや障害回復では重複配送が起こり得るため、注文IDを冪等性キーとしてDBの条件付き書き込みで副作用を一度にする。
3. メッセージは外部副作用のコミット後にだけ削除し、途中障害ならtimeout後に再試行可能にする。
4. 恒久エラーはmaxReceiveCount後にDLQへ隔離し、修正後のredriveで戻すAが、信頼性とスループットを両立する。

4 正解が要件を満たす理由

Aは、長い処理に合わせたvisibility timeoutと動的延長で同時処理を抑えつつ、Standardキューで重複が完全には消えないことを前提に注文IDで副作用を冪等化する。DBコミット前にはメッセージを削除しないため、ワーカー障害時も再実行できる。繰り返し失敗する破損画像はDLQへ分離され、調査・修正後にredriveできるので、正常ジョブの処理能力も守られる。

5・6 各選択肢の検討

不適切な理由と、正解になり得る条件を対にして確認します。

A 正解

可視性の延長、コミット後削除、冪等性、有限回の再試行とDLQを組み合わせ、少なくとも1回配送を安全に扱う。

B 不正解

FIFOの送信時重複排除だけでは、処理がvisibility timeoutを超えた再配送や副作用の再実行を安全にできない。同一グループの毒メッセージは後続も止める。

この選択肢が正解になる条件

厳密な順序が必須で、処理時間内のvisibility timeout、冪等なコンシューマー、DLQを別途構成し、FIFOのスループットで足りる場合。

C 不正解

同じジョブの並列実行を意図的に増やし、外部副作用の競合と計算費用を悪化させる。事後返金は二重課金を防ぐ要件を満たさない。

この選択肢が正解になる条件

処理が副作用を持たず、同じ計算を競争実行して最初の結果だけを採用することが遅延短縮に有効で、重複計算費を許容する場合。

D 不正解

一時障害まで1回でDLQへ送るため再試行性が低く、処理前削除はDB更新前の障害でジョブを永久に失う。保持期間を延ばしても削除済みメッセージは戻らない。

この選択肢が正解になる条件

処理失敗を一度も自動再試行せず必ず人手審査へ回す必要があり、受信直後に別の耐久ストアへジョブを原子的に引き継げる場合。

7 今後使える判断基準

SQSコンシューマーは、処理時間より長いvisibility timeout、必要時の延長、成功後の削除、冪等な副作用をセットで考える。DLQは一時障害を試す余地を残したmaxReceiveCountで毒メッセージだけを隔離する。

8 関連サービスとの比較

Standardキューは高スループットで少なくとも1回配送のため、冪等性が基本となる。FIFOの重複排除は送信重複を抑えるが、コンシューマーの副作用設計を不要にはしない。visibility timeoutは削除ではなく一時的な不可視化である。

QUESTION 007

問題7

PERFORMANCE

本番相当

単一選択

DynamoDBの書き込み分散

IoT分析サービスは、機器イベントをDynamoDBテーブルへ毎秒約4万件保存している。現在のパーティションキーはtenantId、ソートキーはeventTimestamp#eventIdである。最大顧客1社が書き込みの42%を占め、新製品の稼働開始時にその顧客だけWriteKeyRangeThroughputThrottleEventsが急増する。テーブル全体の消費量は設定済み最大値を十分下回り、他顧客は正常である。主要なアクセスパターンは、顧客と1時間の時間範囲を指定した一覧、およびグローバルに一意なeventIdによる1件取得で、いずれも1秒未満が必要である。イベントの順序を顧客全体で厳密に直列化する必要はない。今後10倍の負荷にも対応し、過剰な容量購入や全表Scanを避ける変更はどれか。

- A** 既存キーを維持したままオンデマンドキャパシティモードへ変更し、テーブルの最大スループット上限を10倍にする。adaptive capacityが最大顧客の単一tenantIdへ全テーブル容量を常に移す前提で運用する。
- B** パーティションキーをtenantId#日付#shardIdへ変更し、eventIdのハッシュから固定数のshardIdを決める。時間範囲は該当shardを並列Queryし、eventIdをパーティションキーとするGSIで1件取得する。
- C** 既存テーブルの前段にDAXクラスターをMulti-AZで配置し、書き込みもすべてDAXエンドポイントへ送る。tenantIdは変えず、キャッシュヒット率を高めれば単一論理キーへの書き込みスロットリングも解消する。
- D** 既存テーブルは変更せず、tenantIdをパーティションキーとするGSIを追加して時間範囲の読み取りをGSIへ移す。ベーステーブルとGSIへ容量を均等に割り当て、最大顧客の書き込みを2つの物理経路へ分散する。

ここで答えを決める

正解・解説は次のページから

ANSWER 007

問題7の正解・解説

1 正解 **B**

- B** パーティションキーをtenantId#日付#shardIdへ変更し、eventIdのハッシュから固定数のshardIdを決める。時間範囲は該当shardを並列Queryし、eventIdをパーティションキーとするGSIで1件取得する。

2 問題文の重要な条件

- テーブル全体では余裕があるのに、1社のtenantIdだけが書き込みホットキーになっている
- 顧客・時間範囲の一覧とeventIdによる1件取得を低遅延で維持する
- 顧客全体の厳密な書き込み順序は不要で、10倍成長を見込む
- 全表Scanと単純な過剰プロビジョニングを避ける

3 正解へ至る判断過程

1. テーブル総容量ではなくKeyRangeのスロットリングなので、同じtenantIdへ集中する論理キー設計が原因と判断する。
2. 日付バケットと計算可能なshardIdで1顧客の書き込みを複数パーティションキーへ展開すれば、物理パーティション間で負荷を分散できる。
3. 顧客・時間範囲の検索は対象日と既知のshard数へ限定して並列Queryでき、無制限のScanにはならない。
4. eventIdの高カーディナリティGSIを別アクセスパターンに使うBが、書き込み分散と二種類の読み取りを両立する。

4 正解が要件を満たす理由

Bは、ホットなtenantIdを日付と決定的なwrite shardへ展開し、書き込みを複数のパーティションキーへ均等化する。shard数が既知なので、時間範囲検索は該当日・shardへの並列Queryとして上限を持たせられる。グローバルに一意で高カーディナリティなeventIdをGSIのキーにすれば、1件取得もScanなしで行える。厳密な顧客全体順序が不要という条件が、結果統合を伴うwrite shardingを許している。

5・6 各選択肢の検討

不適切な理由と、正解になり得る条件を対にして確認します。

A 不正解

オンデマンドは総需要への容量管理を軽減するが、単一パーティションキーへ集中する無制限のスループットを保証しない。キー分布を直さず上限だけ増やしてもホットキーは残る。

この選択肢が正解になる条件

アクセスが多数の高カーディナリティキーへ均等に分散し、予測不能なテーブル全体の増減だけが問題である場合。

B 正解

write shardingで最大顧客の書き込みを分散し、限定的な並列QueryとeventIdのGSIで既存アクセスパターンも維持する。

C 不正解

DAXは読み取りキャッシュとしてDB読み取り負荷を減らせるが、同じtenantIdへのベーステーブル書き込み分布は変えない。今回観測された書き込みKeyRangeの問題を解消しない。

この選択肢が正解になる条件

キー分布と書き込み性能は健全で、繰り返しの結果整合性読み取りによる高遅延とRCU消費が主因である場合。

D 不正解

同じtenantIdをGSIのパーティションキーにするとGSIにもホットキーを作り、GSI更新分の書き込みコストも増える。ベーステーブルの書き込みを二経路へ分割する仕組みではない。

この選択肢が正解になる条件

ベーステーブルのキー分布は均等で、別の高カーディナリティ属性によるQueryを追加するため、その属性をキーとしたGSIが必要な場合。

7 今後使える判断基準

DynamoDBでテーブル総容量に余裕があるのに特定キーだけがスロットルされる場合、容量モードより先にパーティションキーのカーディナリティと偏りを見る。write shardingは、読み取り時に探索するshard数を限定できる設計と対にする。

8 関連サービスとの比較

オンデマンドとadaptive capacityは変動や不均衡を吸収するが、悪い単一キー設計を無制限に救済しない。DAXは反復読み取り向けで、書き込み分散ではない。GSIは別アクセスパターンを高速化するが、低カーディナリティキーならGSI自体がホットになる。

QUESTION 008

問題8

COST

本番より難しい

複数選択 (2つ)

中断可能バッチの購入オプション

メディア企業は毎晩、独立した20万本の動画セグメントをEC2で変換する。ジョブはSQSから取得でき、平均8分、最長25分で、処理順序は問わない。午前6時までの4時間に95%以上を完了する必要がある。前日の動画本数によって必要容量は40～800台に変動し、日中の同じ計算環境はほぼ使用しない。アプリケーションは2分ごとに進捗をS3へチェックポイントできるが、現状は1種類のオンデマンドインスタンスだけを使い、ピーク時に800台となるためコストが高い。過去データでは必要容量の15%を常時確保すれば期限を守るための最低処理率を維持でき、残りは中断されても再実行可能である。特定インスタンスタイプの在庫不足にも耐え、Spotの割引を最大限使いながら期限逸脱と重複結果を抑えたい。実施すべき変更を2つ選択してください。

- A** EC2 Auto ScalingのMixedInstancesPolicyで複数の適合ファミリー、サイズ、AZを指定する。必要容量の15%をOn-Demandのベース容量とし、残りをSpotのprice-capacity-optimized配分とCapacity Rebalancingで調達する。
- B** 全容量を過去1か月で最安だった単一インスタンスタイプ、単一AZのSpotへ変更し、Spot上限価格をオンデマンド価格の30%に固定する。中断時は同じプールの容量が戻るまでAuto Scalingの再試行を続ける。
- C** SQSメッセージを作業単位として、結果キーに条件付き書き込みを用いて冪等化する。チェックポイントから再開し、Spot中断通知を受けたワーカーは進捗を保存してメッセージを再実行可能にする。
- D** 全800台をオンデマンドのまま維持し、3年のCompute Savings Plansを夜間ピーク全量に対して購入する。キューの深さにかかわらず毎日4時間起動すれば、Spot中断なしで最小費用になる。

ここで答えを決める

正解・解説は次のページから

ANSWER 008

問題8の正解・解説

1 正解 A・C

- A** EC2 Auto ScalingのMixedInstancesPolicyで複数の適合ファミリー、サイズ、AZを指定する。必要容量の15%をOn-Demandのベース容量とし、残りをSpotのprice-capacity-optimized配分とCapacity Rebalancingで調達する。
- C** SQSメッセージを作業単位として、結果キーに条件付き書き込みを用いて冪等化する。チェックポイントから再開し、Spot中断通知を受けたワーカーは進捗を保存してメッセージを再実行可能にする。

2 問題文の重要な条件

- 最低処理率を支える15%は中断されない容量として必要である
- 残りの独立ジョブは中断と再実行を許容し、Spotを最大限利用できる
- 単一インスタンスタイプや単一AZの在庫不足を避ける
- チェックポイント、SQS、冪等性で期限と重複結果のリスクを抑える

3 正解へ至る判断過程

1. 期限を守る床となる15%はOn-Demand base capacityにし、それを超える弾力部分をSpotへ割り当てる。
2. Spotは複数のインスタンスタイプとAZへ候補を広げ、価格だけでなく利用可能容量を考慮する配分戦略で中断と調達失敗を減らす。
3. インフラ側で中断を減らしても排除はできないため、通知、短いチェックポイント、SQS再実行をアプリケーション側で組み合わせる。
4. 少なくとも1回の再実行で結果が重複しないよう条件付き書き込みで冪等化するAとCが一体の解になる。

4 正解が要件を満たす理由

Aは、期限達成に必要な床をOn-Demandで守りながら、追加容量を複数プールのSpotから調達してコストを下げる。price-capacity-optimizedとCapacity Rebalancingは、最安単一プールへの集中より調達成功率と継続性を高める。Cは、中断を通常事象としてチェックポイントから再開し、SQSメッセージを再実行可能に保つ。結果の条件付き書き込みにより再配送・再実行時の重複も防ぐため、両方でコストと期限を両立する。

5・6 各選択肢の検討

不適切な理由と、正解になり得る条件を対にして確認します。

A 正解

On-Demandの最低容量と多様なSpotプールを1つのAuto Scalingグループで混在させ、可用性を保ちながら割引対象を最大化する。

B 不正解

最安の単一プールへ集中すると容量不足や一斉中断の影響が大きく、15%の最低処理率も保証しない。低い上限価格は必要容量を取得できない要因になる。

この選択肢が正解になる条件

完了期限がなく、処理を何日でも延期でき、単一タイプ専用バイナリのため他のインスタンスプールを利用できない場合。

C 正解

Spot中断後の再開量を小さくし、SQS再実行と冪等な結果確定により、少なくとも1回処理でも安全に進捗を回復できる。

D 不正解

中断リスクは小さいが、停止許容の75%以上の弾力部分までオンデマンドにするためSpot節約を得られない。夜間ピーク全量の3年確約は利用変化や未使用時間のリスクも大きい。

この選択肢が正解になる条件

処理を一度も中断できず全容量が数年間ほぼ常時稼働し、確約した時間当たり利用額を昼夜とも確実に消費する場合。

7 今後使える判断基準

Spot適性は中断可能だけでなく、最低保証容量、候補プールの多様性、チェックポイント、再実行、冪等性で判断する。期限のあるバッチはOn-Demandで床を作り、弾力部分だけを複数プールのSpotへ載せる。

8 関連サービスとの比較

On-Demandは中断されない基礎容量に適し、Spotはフォールトトレラントな追加容量に適する。MixedInstancesPolicyは両購入オプションと複数タイプを混在できる。Savings Plansは利用料金の割引であり、Spot中断対策や容量確保の仕組みではない。

QUESTION 009

問題9

SECURE

複合難問

単一選択

PIIの直接アップロード統制

保険会社は、顧客のブラウザから100MB～4GBの診断画像をS3へ直接アップロードする新機能を作る。APIサーバーを大容量データの中継点にせず、顧客へAWS認証情報を一切渡さない。アップロード権限は申請後15分、サーバーが採番した顧客固有オブジェクトキー1個だけに限定する。同名キーは事前に存在せず、アップロード完了後の非同期イベントでマルウェア検査を開始する。バケットとオブジェクトは公開せず、通信はTLS、保存時は監査部門所有の特定カスタマーマネージドKMSキーによるSSE-KMSを必須とし、SSE-S3や別KMSキーを明示したPUTは拒否する。署名URLが漏えいした場合の有効期間と操作範囲も最小化する。URLを入手した顧客が他のキーを上書きできないようにし、KMS利用とS3データイベントを監査したい。最も適切な構成はどれか。

- A** バケットのデフォルト暗号化を対象KMSキーへ設定し、全顧客で共通のプレフィックスに対する24時間有効なpresigned URLを発行する。暗号化ヘッダーの有無は検査せず、S3 Block Public Accessだけで誤った暗号化方式も防ぐ。
- B** Amazon Cognito Identity Poolから各ブラウザへ一時AWS認証情報を渡し、バケット全体のPutObjectを許可する。ブラウザが任意キーと任意の暗号化ヘッダーを指定し、アップロード後にLambdaで不適切なオブジェクトを削除する。
- C** 限定IAMロールで、採番キーとSSE-KMS方式・対象キーの各ヘッダーを署名に含む15分のpresigned PUT URLを生成する。bucket policyでTLS、SSE-KMS、対象キーを強制し、Block Public AccessとKMSキーポリシー、監査ログを設定する。
- D** ブラウザからAPI Gatewayへ画像をBase64で送信し、Lambdaが対象KMSキーを指定してS3へPutObjectする。APIの認証を15分で失効させれば、4GBのファイルも単一リクエストで安全に中継できる。

ここで答えを決める

正解・解説は次のページから

ANSWER 009

問題9の正解・解説

1 正解 C

- C** 限定IAMロールで、採番キーとSSE-KMS方式・対象キーの各ヘッダーを署名に含む15分のpresigned PUT URLを生成する。bucket policyでTLS、SSE-KMS、対象キーを強制し、Block Public AccessとKMSキーポリシー、監査ログを設定する。

2 問題文の重要な条件

- 100MB～4GBをAPIサーバーで中継せずブラウザからS3へ直接送る
- AWS認証情報を顧客へ渡さず、15分・採番済み1キーだけを許可する
- TLSと特定カスタマーマネージドKMSキーのSSE-KMSを強制する
- 公開、別キー上書き、SSE-S3や別KMSキーへの迂回を防ぎ監査する

3 正解へ至る判断過程

1. presigned URLは、署名したIAMプリンシパルの権限を特定操作・特定キー・期限へ委任でき、ブラウザへAWS認証情報を配布しない。
2. 暗号化方式とKMSキーのヘッダーを署名時に固定すれば、クライアントが変更した要求は署名と一致しない。
3. デフォルト暗号化だけでは、クライアントが別方式を明示する要求を拒否する統制にならないため、bucket policyの明示的DenyでTLS、方式、キーを強制する。
4. 署名ロールとKMSキーポリシーの最小権限、Block Public Access、CloudTrail監査まで備えるCが全条件を満たす。

4 正解が要件を満たす理由

Cでは、バックエンドが採番した単一オブジェクトキーに対する短命のpresigned PUT URLだけを返すため、認証情報配布や他キーへの権限拡大を避けられる。SSE-KMS方式とKMSキーIDを署名対象ヘッダーに含め、bucket policyでも安全な通信、SSE-KMS、指定キー以外を明示的に拒否することで迂回を防ぐ。署名ロールにはPutObjectと必要なKMSデータキー権限だけを与え、S3データイベントとKMSイベントを記録できる。

5・6 各選択肢の検討

不適切な理由と、正解になり得る条件を対にして確認します。

A 不正解

デフォルト暗号化はヘッダー未指定時の方式を決めるが、別方式を明示したPUTを拒否する統制ではない。共通プレフィックスと長い有効期限も1キー・15分の範囲を超える。

この選択肢が正解になる条件

アップロード元が信頼済みで、任意のキーへの書き込みと長い期限を許容し、暗号化はデフォルト方式になればよく別方式を明示的に拒否する必要がない場合。

B 不正解

一時的でもAWS認証情報をブラウザへ渡す明示禁止に反し、バケット全体へのPutObjectは権限が広すぎる。事後削除では公開や誤暗号化の保存を予防できない。

この選択肢が正解になる条件

クライアントへ一時AWS認証情報を渡せて、IAM条件で各ユーザーのプレフィックスと暗号化ヘッダーを厳密に制限し、複数オブジェクト操作が必要な場合。

C 正解

短命・単一キーの署名と、要求ヘッダー・bucket policy・KMS権限の多層統制により、直接アップロードの性能とPII保護を両立する。

D 不正解

API GatewayとLambdaで4GBの単一ペイロードを中継する設計はサービスのペイロード上限を超え、API層を大容量データ経路にしない条件にも反する。

この選択肢が正解になる条件

ファイルがAPI GatewayとLambdaのペイロード上限内の小容量で、マルウェア検査などの理由から同期的なアプリケーション中継が必須の場合。

7 今後使える判断基準

S3への外部直接アップロードでは、presigned URLの期限・HTTP操作・オブジェクトキー・必須ヘッダーを固定し、bucket policyの明示的DenyでTLSと暗号化方式・キーを強制する。デフォルト暗号化だけを強制策と誤認しない。

8 関連サービスとの比較

presigned URLは署名者の権限を特定要求へ期限付きで委任し、Cognitoはクライアントへ一時AWS認証情報を渡して複数APIを許可できる。SSE-S3はS3管理キー、SSE-KMSはキーポリシーと利用監査を持つKMSキーを使う。

QUESTION 010

問題10

RESILIENT

本番相当

単一選択

イベントの独立ファンアウト

旅行会社の予約サービスは、予約作成・変更・取消イベントを既存のEventBridgeカスタムイベントバスへ発行している。通常は毎秒2万件、セール時は毎秒5万件となり、JSON本文の複数フィールドを組み合わせた振り分けが必要である。決済、顧客通知、分析の3チームが、それぞれ異なるイベント内容を選別して処理する。各コンシューマーは独立してデプロイされ、最大6時間停止することがあり、復旧後は自分宛ての未処理イベントを速度制御しながら消化したい。1チームの障害が他チームへ波及してはならず、恒久的に失敗するメッセージはチーム別に隔離する。厳密な順序は不要である。さらに監査担当は過去7日分を保持し、ルール修正後に該当イベントを再送したい。将来は別AWSアカウントのコンシューマーも追加する予定で、プロデューサーは変更したくない。最も適切な構成はどれか。

- A** イベントパターンごとのEventBridgeルールから3つのLambdaを直接呼び出し、全ターゲットで1つのSQS DLQを共有する。コンシューマー停止中はLambda予約済み同時実行数を0にし、EventBridgeの標準再試行だけを未処理バッファとする。
- B** 全イベントを1本のSQS Standardキューへ送信し、3チームのワーカーが同じキューをポーリングする。各ワーカーが不要なイベントを受信したら削除し、必要なチームが取得するまでvisibility timeoutで競合を調整する。
- C** EventBridgeから1つのSNS Standardトピックへ全イベントを転送し、3本のSQSキューを購読させてフィルタする。SNSとSQSの保持期間を7日にすれば、後から作成したルールへ過去イベントを任意に再生できる。
- D** チーム別のEventBridgeルールで内容を照合し、それぞれ専用SQS Standardキューへ送る。各キューにコンシューマー用DLQを設け、ターゲット配送にもDLQと再試行を設定する。イベントバスのアーカイブを7日保持してリプレイする。

ここで答えを決める

正解・解説は次のページから

ANSWER 010

問題10の正解・解説

1 正解 **D**

- D** チーム別のEventBridgeルールで内容を照合し、それぞれ専用SQS Standardキューへ送る。各キューにコンシューマー用DLQを設け、ターゲット配送にもDLQと再試行を設定する。イベントバスのアーカイブを7日保持してリプレイする。

2 問題文の重要な条件

- 既存のEventBridgeバスを変えず、内容に応じて複数チームへファンアウトする
- 各コンシューマーは6時間停止しても独立した未処理分を保持し、速度制御して再開する
- 配送失敗と処理失敗を失わず、チーム別に隔離する
- 7日分の過去イベントをルール修正後にリプレイし、将来はクロスアカウントへ拡張する

3 正解へ至る判断過程

1. ファンアウトでは、同じイベントを必要な各コンシューマーへ複製する必要がある、競合コンシューマーが1本のキューを奪い合う構成は除外する。
2. EventBridgeルールは内容ベースで各チーム宛てを選別でき、プロデューサーと追加先を分離できる。
3. チームごとのSQSキューをターゲットにすると、停止中も耐久的に蓄積し、各チームが独自速度でポーリングできる。配送段階と処理段階のDLQを分ける。
4. EventBridgeアーカイブとリプレイが過去イベントの再送要件に直接対応するためDを選ぶ。

4 正解が要件を満たす理由

Dは、EventBridgeのルールをチーム別に分け、各ターゲットを専用SQSキューにすることで、ファンアウト、内容ベースルーティング、障害分離、バックプレッシャーを同時に実現する。EventBridgeからSQSへの配送不能はターゲットDLQ、SQSから取得後の恒久的な処理失敗は各キューのredrive DLQで別々に追跡できる。イベントバスのアーカイブは7日分を保持してルール修正後にリプレイでき、別アカウント追加でもプロデューサーを変えずに済む。

5・6 各選択肢の検討

不適切な理由と、正解になり得る条件を対にして確認します。

A 不正解

直接呼び出しは短い障害への再試行はできるが、チームが6時間分の滞留を独自速度で消化する耐久キューを持たない。共有DLQでは障害所有者と再投入先も分離しにくい。

この選択肢が正解になる条件

各処理が短時間で常時利用可能で、コンシューマー側のバックプレッシャーや長時間停止後の速度制御が不要な単純なリアルタイム処理の場合。

B 不正解

1本のSQSキューではメッセージは1コンシューマーに処理される競合モデルとなり、3チーム全員へのファンアウトにならない。不要として削除すれば必要なチームへ届かない。

この選択肢が正解になる条件

同じ種類のワーカー群のうち任意の1台だけが各タスクを処理すればよく、購読者ごとの複製が不要な作業キューの場合。

C 不正解

SNSとチーム別SQSによるファンアウト自体は有効だが、保持期間を設定したSNS/SQSが、新しいルールへ過去イベントを任意に再生するEventBridgeアーカイブの代わりにはならない。

この選択肢が正解になる条件

過去イベントのアーカイブ・リプレイと複雑なイベントバスルーティングが不要で、現在のメッセージを購読フィルタ付きで複数SQSへ即時配信するだけの場合。

D 正解

EventBridgeのルーティングとアーカイブ、チーム別SQSの耐久バッファ、段階別DLQを組み合わせ、独立停止・再開・再生を満たす。

7 今後使える判断基準

イベント通知と作業の耐久性を分け、EventBridgeで内容ベースのルーティング、購読者ごとのSQSで停止吸収と速度制御を行う。再試行不能時は配送DLQ、処理不能時はコンシューマーDLQ、過去再送はアーカイブで扱う。

8 関連サービスとの比較

EventBridgeは高度なイベントパターン、クロスアカウント経路、アーカイブとリプレイに向く。SNSは低遅延のpub/subファンアウト、SQSはpull型の耐久バッファとバックプレッシャーに向く。1本のSQSは複数購読者へのブロードキャストではない。

QUESTION 011

問題11

PERFORMANCE

本番より難しい

単一選択

共有セッションとキャッシュの可用性

チケット販売サイトは3つのAZにまたがるEC2 Auto Scalingグループで動作している。現在、ログインセッションと商品検索結果を各EC2のメモリに保存しているため、スケールインやAZ障害でログアウトが多発し、同じ検索がAuroraへ集中する。ピークは毎秒15万リクエスト、共有データは現在220GBで1年以内に500GBへ増える。セッション更新は1つのsessionId内で原子的に行い、TTLで30分後に失効させたい。キャッシュは再生成可能で、フェイルオーバー時に直近数秒のセッション更新を失うことは許容するが、ノードまたはAZ障害でセッション格納領域全体を失ってはならず、2分以内に自動復旧したい。既存コードはRedis互換クライアントを導入できる。最小の運用変更で性能、水平拡張、可用性を満たす構成はどれか。

- A** ElastiCache for Redis OSSをクラスターモード有効の複数シャードで作成し、各シャードに別AZのレプリカを置いてMulti-AZ自動フェイルオーバーを有効化する。セッションと検索結果にTTLを設定し、設定エンドポイントを利用する。
- B** ElastiCache for Memcachedの多数ノードを3つのAZへ分散し、クライアント側consistent hashingでセッションと検索結果を配置する。レプリカは作らず、障害ノードのキーはTTLまで他ノードから自動復元される前提とする。
- C** ElastiCache for Redis OSSのクラスターモード無効、リードレプリカなしの単一大容量ノードへ全データを集約する。毎日スナップショットを取得し、障害時は最新バックアップから新ノードを手動復元する。
- D** セッションを各EC2のローカルメモリに残し、ALBのスティッキーセッションを30分に設定する。商品検索結果だけを各AZの独立したMemcachedへ置き、AZ障害時はRoute 53で正常AZのALBへ切り替える。

ここで答えを決める

正解・解説は次のページから

ANSWER 011

問題11の正解・解説

1 正解 **A**

- A** ElastiCache for Redis OSSをクラスターモード有効の複数シャードで作成し、各シャードに別AZのレプリカを置いてMulti-AZ自動フェイルオーバーを有効化する。セッションと検索結果にTTLを設定し、設定エンドポイントを利用する。

2 問題文の重要な条件

- 共有データが220GBから500GBへ増え、毎秒15万リクエストを水平分散する
- sessionId単位の原子的更新と30分TTLが必要である
- ノードまたはAZ障害でも全セッションを失わず、2分以内に自動復旧する
- 直近数秒の損失は許容し、再生成可能な検索キャッシュと同じRedis互換層を使える

3 正解へ至る判断過程

1. ローカルメモリから共有キャッシュへ移すだけでなく、将来容量を単一ノードへ閉じ込めないためデータ分割が必要である。
2. Redis OSSのクラスターモードは複数シャードへキーを分散し、sessionIdごとの単一キー操作なら原子性とTTLを維持しやすい。
3. 各シャードに別AZのレプリカを置きMulti-AZ自動フェイルオーバーを使えば、一次ノードやAZ障害時にレプリカを昇格できる。
4. Memcachedは単純な水平キャッシュには強いがノード間レプリケーションと自動フェイルオーバーを持たないため、セッション要件にはAが適する。

4 正解が要件を満たす理由

Aは、Redis OSSのシャーディングで220GBから500GBへの容量と高いリクエスト率を複数ノードへ分散する。各セッションを単一キーとして扱えば、原子的更新とTTLを利用できる。各シャードのレプリカを別AZに置きMulti-AZ自動フェイルオーバーを有効にすることで、障害時にセッション領域全体を失わず短時間で昇格できる。検索結果はcache-asideで再生成でき、同じ共有層へ置くことでAurora負荷も下げられる。

5・6 各選択肢の検討

不適切な理由と、正解になり得る条件を対にして確認します。

A 正解

Redis互換の原子的操作とTTL、クラスターモードの水平分割、レプリカとMulti-AZ自動フェイルオーバーを一つの構成で満たす。

B 不正解

Memcachedはマルチスレッドで単純キャッシュの水平拡張に強いが、ノード間レプリケーションを提供しない。障害ノード上のセッションは自動的に他ノードから復元されない。

この選択肢が正解になる条件

格納物がすべて即時再生成可能で、ノード障害時のキャッシュミスと一部キー消失を許容し、単純なオブジェクトキャッシュのスループットを優先する場合。

C 不正解

Redisの機能は使えるが、単一ノードが容量・性能・可用性の障害点になる。日次スナップショットからの手動復元は2分RTOも直近数秒のRPOも満たさない。

この選択肢が正解になる条件

データ量と負荷が単一ノードに十分収まり、キャッシュ全損を許容でき、レプリカ費用を避ける開発・検証環境の場合。

D 不正解

スティッキーセッションはEC2やAZが失われたときローカル状態を移送せず、Auto Scalingのスケールインでもログアウトを防げない。AZ別キャッシュも共有性と自動フェイルオーバーを欠く。

この選択肢が正解になる条件

セッションが不要なステートレス処理、またはノード消失時の全ログアウトを許容する小規模環境で、短時間の接続親和性だけが必要な場合。

7 今後使える判断基準

ElastiCacheのエンジン選択では、単純な再生成可能キャッシュならMemcached、レプリケーション、自動フェイルオーバー、複雑・原子的データ操作が必要ならRedis OSSを優先する。容量が単一ノードを超えて成長するならRedisのクラスターモードも検討する。

8 関連サービスとの比較

Memcachedはマルチスレッドの単純オブジェクトキャッシュとノード追加・削除に向くが、レプリケーションやバックアップを持たない。Redis OSSはTTL、豊富なデータ型、レプリカ、Multi-AZ自動フェイルオーバーを利用でき、クラスターモードでシャーディングできる。

QUESTION 012

問題12

COST

複合難問

複数選択 (2つ)

負荷形状に応じたECS起動タイプ

メディア企業は、同じコンテナイメージで2つのECSワークロードを運用する。サービスXは画像変換APIで、平日は無通信の時間が多いがキャンペーン時だけ数分で0から80タスクへ増え、事前に時刻を予測できない。インスタンス管理担当者はいない。サービスYは動画トランスコード制御で、合計256 vCPUを24時間365日、平均75%以上利用し、今後3年間の基礎負荷はほぼ一定である。Yは複数インスタンスタイプで実行でき、タスクの再配置を許容するが、OSパッチやクラスター容量管理を担当できる小規模SREチームはいる。直近12か月の請求では、Xの待機容量とYのタスク単価が主要費用である。コンテナ改修はせず、購入方式と起動タイプはサービスごとに变えてよい。両サービスの可用性を複数AZで維持しつつ、ワークロードごとの総コストを最小化したい。採用すべき方針を2つ選択してください。

- A** サービスYはECS on EC2へ配置し、複数AZのAuto Scalingグループをcapacity providerで管理する。基礎容量を右サイジングし、定常利用分には適切なSavings Plansを適用する。
- B** サービスXとYをすべてFargate Spotだけで実行する。基礎タスクも中断を前提にし、Savings PlansやOn-Demand容量は使用しない。
- C** サービスXはECS on EC2へ配置し、キャンペーン最大時の80タスクを収容するインスタンスを3AZで常時起動する。未使用時間の容量は可用性予備として保持する。
- D** サービスXはFargateで実行し、ECS Service Auto Scalingで需要に応じてタスク数を増減する。通常は最小タスクだけを維持し、EC2ホストの調達・パッチ適用を不要にして運用工数も抑える。

ここで答えを決める

正解・解説は次のページから

ANSWER 012

問題12の正解・解説

1 正解 A・D

- A** サービスYはECS on EC2へ配置し、複数AZのAuto Scalingグループをcapacity providerで管理する。基礎容量を右サイジングし、定常利用分には適切なSavings Plansを適用する。
- D** サービスXはFargateで実行し、ECS Service Auto Scalingで需要に応じてタスク数を増減する。通常は最小タスクだけを維持し、EC2ホストの調達・パッチ適用を不要にして運用工数も抑える。

2 問題文の重要な条件

- Xは予測不能な急増と長いアイドルがあり、インフラ担当者がいない
- Yは256 vCPUを3年間ほぼ一定かつ高利用率で使う
- Multi-AZ可用性を保ち、コンピューティング費用と運用人件費の双方を評価する

3 正解へ至る判断過程

1. 起動タイプはコンテナ機能ではなく、負荷の継続性、ホスト運用能力、購入割引の柔軟性で分ける。
2. Xはアイドル容量を抱えず急増できるFargateの運用価値が高く、Yは高い定常利用率のためEC2を詰めてコミットメント割引を使う余地が大きい。
3. 中断可能容量だけで常時サービスを構成せず、最大需要分のEC2をX用に常時保持する過剰プロビジョニングも避ける。

4 正解が要件を満たす理由

AはYの予測可能な高利用率を、ECS on EC2のタスク集約、capacity providerによる容量連動、コミットメント割引へ結び付けるため、管理負荷を受け入れる代わりに定常大規模計算の単価を下げられる。DはXの予測不能な短時間バーストでホストの待機容量をなくし、管理者不在という制約も満たす。

5・6 各選択肢の検討

不適切な理由と、正解になり得る条件を対にして確認します。

A 正解

長期・高稼働・大規模なYではEC2容量を効率よく詰め、割引を適用する効果がホスト管理負荷を上回る。

B 不正解

Spotは中断可能な追加処理には有効だが、両サービスの基礎容量までSpotだけにすると容量不足・中断時の可用性を満たせない。Yの安定利用にはコミットメント割引の方が予測しやすい。

この選択肢が正解になる条件

全タスクがチェックポイント可能なバッチで、中断や開始遅延を許容し、期限内に再試行できるなら Fargate Spotのみ、または大半をSpotにする構成が正解になり得る。

C 不正解

Xの最大80タスク分を常時EC2で保持すると長い無通信時間にも支払いが発生し、担当者不在なのに OS・容量運用も増える。可用性予備と最大需要の全量常駐は同義ではない。

この選択肢が正解になる条件

Xの負荷が常に高く予測可能で、起動遅延を一切許容せず、既存SREがEC2を管理できるなら正解になり得る。

D 正解

Xの需要時だけタスク課金し、急な拡張とホスト保守をAWSへ委ねられるため、負荷形状と運用制約に合う。

7 今後使える判断基準

ECSの起動タイプは、短く変動が大きく運用要員が少ない負荷ならFargate、長期・高稼働・大規模でホストを高密度利用できる負荷ならECS on EC2を基準に総コストで比較する。

8 関連サービスとの比較

Fargateはタスク単位課金とホスト管理不要が強みで、変動負荷に向く。ECS on EC2はインスタンス管理が必要だが、定常大規模負荷でbinpack、複数購入方式、特殊インスタンスを活用しやすい。Fargate SpotやEC2 Spotは中断許容部分へ限定して混在させる。

QUESTION 013

問題13

SECURE

本番相当

単一選択

公開APIのDDoSとSQLインジェクション対策

スタートアップは、CloudFrontを入口としてALB上の公開REST APIを世界向けに提供している。最近、単一IPからの大量リクエストと、クエリ文字列にSQL断片を含む試行が増えた。ネットワーク層・トランスポート層の一般的なDDoS攻撃に対する保護に加え、HTTPレイヤーでSQLインジェクションや異常なリクエストレートをオリジン到達前に遮断したい。専任の24時間セキュリティチームはなく、一般的な攻撃パターンの更新を自社で追従したくない。APIのURLとTLS終端は変更できず、正常なフラッシュセールの急増は許可しながら、送信元単位の濫用を可視化する必要もある。現時点ではDDoS対応チームへの直接支援や攻撃時のコスト保護は必須でなく、月額固定費も抑えたい。誤検知はログを見て段階調整できる。最も適切な構成はどれか。

- A** ALBのセキュリティグループで既知の攻撃元IPを拒否し、AWS Shield AdvancedをALBだけに有効化する。CloudFrontでは追加ルールを設定しない。
- B** CloudFrontにAWS WAF Web ACLを関連付け、AWS Managed RulesのSQLi対策とrate-based ruleを設定する。自動適用されるShield StandardとCloudFrontのエッジ保護も利用する。
- C** CloudFrontを外し、NLBをAPIの入口にする。NACLでHTTPリクエスト本文とクエリ文字列を検査し、送信元ごとのレート制限を設定する。
- D** Amazon GuardDutyとAmazon Inspectorを有効化し、検出結果をEventBridgeでLambdaへ送り、悪性HTTPリクエストをリアルタイムに識別してCloudFrontで遮断する。

ここで答えを決める

正解・解説は次のページから

ANSWER 013

問題13の正解・解説

1 正解 **B**

- B** CloudFrontにAWS WAF Web ACLを関連付け、AWS Managed RulesのSQLi対策とrate-based ruleを設定する。自動適用されるShield StandardとCloudFrontのエッジ保護も利用する。

2 問題文の重要な条件

- CloudFront公開APIをオリジン到達前に保護する
- 一般的DDoSに加え、SQLiと送信元レートをHTTPレイヤーで制御する
- 管理ルールを使って運用負荷と固定費を抑える

3 正解へ至る判断過程

1. DDoSのネットワーク・トランスポート層と、SQLiなどのアプリケーション層を別の制御として扱う。
2. Shield StandardはCloudFrontを含む対象サービスに一般的DDoS保護を自動提供し、WAFはWebリクエスト内容とレートを判定できる。
3. Managed RulesをCloudFrontで評価すればオリジン負荷を減らし、Shield Advancedの追加機能が不要という費用条件にも合う。

4 正解が要件を満たす理由

BはShield Standardで一般的なL3/L4 DDoSを追加固定費なしで保護し、AWS WAFの管理ルールとrate-based ruleでSQLiおよびアプリケーション層の大量要求をエッジで遮断する。CloudFrontを維持するためグローバルな吸収能力とオリジン保護も利用でき、専任チーム不在の条件に合う。

5・6 各選択肢の検討

不適切な理由と、正解になり得る条件を対にして確認します。

A 不正解

セキュリティグループは許可ルール中心でHTTP内容や送信元ごとのレートを検査しない。Shield AdvancedだけでもSQLi検査の代替にはならず、CloudFront入口で遮断しないため要件と費用が合わない。

この選択肢が正解になる条件

高度なDDoS可視化、SRT支援、コスト保護、プロアクティブ対応が必要な重要サービスで、WAFも併用するならShield Advancedが正解の一部になり得る。

B 正解

L3/L4の標準DDoS保護とL7のSQLi・レート制御を、それぞれ適切なサービスでCloudFront入口へ適用する。

C 不正解

NACLはステートレスなL3/L4制御で、HTTP本文・クエリ文字列のSQLiやリクエストレートを解釈できない。NLBへの変更もHTTP保護能力を高めない。

この選択肢が正解になる条件

非HTTPのTCP/UDPサービスで静的IPや超低遅延が必要で、IP・ポート単位の制御だけが要件ならNLBとNACLが構成要素になり得る。

D 不正解

GuardDutyは脅威検出、Inspectorは脆弱性管理であり、CloudFrontの個々のHTTP要求をインラインでSQLi判定・レート遮断するサービスではない。

この選択肢が正解になる条件

要件がAWS環境内の不審活動検出やEC2/ECR/Lambdaの脆弱性評価と事後対応自動化なら正解になり得る。

7 今後使える判断基準

Web攻撃問題では、L3/L4 DDoSはShield、HTTP内容・ボット・レートはWAF、配信とエッジ吸収はCloudFrontと役割分担する。高度支援やコスト保護が明示されたときだけShield Advancedを選ぶ。

8 関連サービスとの比較

Shield Standardは一般的DDoS保護を自動提供し、Shield Advancedは追加可視化、SRT、コスト保護などを提供する。AWS WAFはSQLi・XSS・レートなどL7を検査する。セキュリティグループ/NACLはIP・ポート中心、GuardDuty/Inspectorは検出・評価でインラインWAFではない。

QUESTION 014

問題14

RESILIENT

本番より難しい

単一選択

FargateサービスのAZ障害耐性

ECサイトはALBの背後でECS Fargateサービスを実行している。現在、サービスにはdesired count 2が設定され、1つのプライベートサブネットだけが指定されている。タスク障害時はECSが置換するが、そのAZで障害が起きるとAPI全体が停止した。新設計では3AZを使用し、1AZを失っても通常ピークの全トラフィックを処理しなければならない。負荷試験では通常ピークに4タスク分が必要で、新しいタスクの起動とヘルス判定には最大3分かかる。障害直後に容量不足となる設計は許されない。デプロイ中も利用可能な正常タスク数を維持し、起動後にALBヘルスチェックを通らないリビジョンは自動的に失敗として戻したい。EC2ホストの運用は増やさず、アプリケーションはステートレスのままとする。最も適切な構成はどれか。

- A** Fargateサービスは1AZのままdesired countを6へ増やし、ALBだけ3AZを有効にする。タスク数が増ければAZ障害時にもALBが同じタスクへ再接続できる。
- B** 各AZに独立したECSクラスターとALBを1つずつ作り、Route 53の単一レコードでランダムに返す。デプロイ失敗は運用担当者が各クラスターで個別に判断して切り戻す。
- C** ALBとFargateサービスに3AZのサブネットを指定し、必要容量にAZ障害分の余裕を加えたタスク数をAZ間で分散・再均衡する。ECS deployment circuit breakerとロールバック、適切なminimumHealthyPercentを設定する。
- D** Fargateを単一の大型EC2インスタンス上のECSへ移行し、Auto Scalingのdesired capacityを1にする。インスタンス再起動でタスク障害とAZ障害の両方へ対応する。

ここで答えを決める

正解・解説は次のページから

ANSWER 014

問題14の正解・解説

1 正解 C

- C** ALBとFargateサービスに3AZのサブネットを指定し、必要容量にAZ障害分の余裕を加えたタスク数をAZ間で分散・再均衡する。ECS deployment circuit breakerとロールバック、適切なminimumHealthyPercentを設定する。

2 問題文の重要な条件

- 3AZを使用し、1AZ喪失後もピークトラフィック全量进行处理する
- デプロイ中の正常容量を保ち、ヘルスチェック失敗を自動ロールバックする
- Fargateとステートレス設計を維持し、EC2運用を増やさない

3 正解へ至る判断過程

1. ロードバランサーだけでなく、実際のFargateタスクを複数AZのサブネットへ分散させる必要がある。
2. AZ障害後にもピーク全量进行处理するには、通常時からN+1相当のタスク余力を持ち、ECSの再均衡・置換を利用する。
3. デプロイ可用性はminimumHealthyPercentなど、失敗判定はcircuit breakerとALBヘルスにより自動化する。

4 正解が要件を満たす理由

CはALBとタスクの両方を3AZへ配置し、AZ障害分を含む容量を事前確保するため、1AZ喪失直後にも残存タスクでピーク进行处理できる。ECSサービスがdesired countを維持・再均衡し、circuit breakerとロールバックが不健全な新リビジョンの長時間展開を防ぐ。Fargateの運用モデルも保つ。

5・6 各選択肢の検討

不適切な理由と、正解になり得る条件を対にして確認します。

A 不正解

ALBをMulti-AZにしても、全ターゲットが1AZならそのAZ障害で接続先がなくなる。タスク数だけ増やしても障害ドメインの分離にならない。

この選択肢が正解になる条件

要件がタスクプロセス障害だけへの耐性で、AZ障害を明示的に対象外とする開発環境なら単一AZでの複数タスクが正解になり得る。

B 不正解

独立クラスターとALBを3組作ると設定・デプロイ・DNS管理が増え、単一ECSサービスで実現できるAZ分散より運用負荷が高い。ランダムDNSは正常性と容量を十分に扱わない。

この選択肢が正解になる条件

AZごとに厳密な障害・デプロイ境界を持たせる特殊なセルアーキテクチャが必要で、自動化基盤を用意できるなら正解になり得る。

C 正解

タスク配置、容量余力、ロードバランシング、置換、デプロイ失敗処理を一体で設計し、全要件を満たす。

D 不正解

単一EC2・単一AZが新たな単一障害点となり、Fargateのホスト管理不要という制約にも反する。インスタンス再起動はAZ障害を解決しない。

この選択肢が正解になる条件

特殊なホスト機能が必要で、複数AZ・複数EC2のAuto Scalingグループとcapacity providerを構成するならECS on EC2が正解になり得る。

7 今後使える判断基準

ECS高可用性では、ALBのAZ数だけでなくタスクのサブネット、desired count、AZ喪失後の必要容量を確認する。デプロイ障害は正常率設定、ヘルスチェック、circuit breaker/rollbackまで含める。

8 関連サービスとの比較

Fargateはホストを意識せず複数AZへタスク配置でき、ECSサービスがdesired countを維持する。ALBは正常ターゲットへHTTPを分散するが、ターゲット自体のAZ分散は代行しない。ECS on EC2ではさらにEC2 Auto Scaling側のAZ容量設計が必要になる。

QUESTION 015

問題15

PERFORMANCE

複合難問

単一選択

非HTTP通信の静的IPロードバランシング

産業IoT企業は、世界中の装置から独自バイナリプロトコルをTCP 9000番で受信し、同じエンドポイントでUDP 9001番のテレメトリも受ける。接続は長時間継続し、レイヤー7のHTTPルーティングは不要である。顧客工場のファイアウォールでは接続先を少数の固定IPv4アドレスで許可リスト登録する必要がある。バックエンドは3AZのEC2で、毎秒数百万パケット時にも低遅延で処理したい。TLSはアプリケーション側で終端し、送信元IPを利用した監査を維持する。既存装置はHTTPへ変換できず、プロトコルを中継する独自プロキシの保守も避けたい。AZ障害時は正常ターゲットへの新規接続を継続する必要がある。DNS名だけに依存する動的IP変更は受け入れられない。性能試験も実施する。最も適切なロードバランサー構成はどれか。

- A** Application Load Balancerを3AZで有効化し、TCP 9000とUDP 9001のリスナーを追加する。固定IPはALBノードの現在のアドレスを定期取得して顧客へ配布する。
- B** Classic Load Balancerを使用し、TCPリスナーとHTTPヘルスチェックを構成する。UDPは同じリスナーへカプセル化し、Elastic IPをCLBへ関連付ける。
- C** Gateway Load Balancerをインターネット向けに公開し、GENEVEでIoT装置からのTCP/UDPを直接受信する。各AZのGWLBエンドポイントにElastic IPを付与する。
- D** Network Load Balancerを3AZで構成し、TCPとUDPのリスナーおよび適切なターゲットグループを作る。各AZのNLBノードにElastic IPを割り当て、クライアントIPを保持する。

ここで答えを決める

正解・解説は次のページから

ANSWER 015

問題15の正解・解説

1 正解 **D**

- D** Network Load Balancerを3AZで構成し、TCPとUDPのリスナーおよび適切なターゲットグループを作る。各AZのNLBノードにElastic IPを割り当て、クライアントIPを保持する。

2 問題文の重要な条件

- HTTPではないTCPとUDPの両方を同じ公開基盤で処理する
- 顧客の許可リスト用に少数の静的IPv4が必要
- 高パケットレート、低遅延、長時間接続、送信元IP保持が重要

3 正解へ至る判断過程

1. L7のパス・ホストルーティングが不要でTCP/UDPが必要なため、ALBではなくL4ロードバランサーを選ぶ。
2. NLBはTCP/UDPリスナー、高スループット・低遅延、クライアントIP保持に対応し、インターネット向けNLBではAZごとにElastic IPを割り当てられる。
3. GWLBはセキュリティアプライアンスの透過的な挿入、ALBはHTTP/HTTPS、CLBは新規のUDP・静的IP要件に合わない。

4 正解が要件を満たす理由

Dは必要なTCPとUDPをネイティブに受け、高パケットレートのL4処理と長時間接続に適する。AZごとのElastic IPにより顧客は固定IPv4を許可リスト化でき、3AZの正常なEC2へ分散できる。クライアントIP保持とアプリ側TLS終端も要件に矛盾しない。

5・6 各選択肢の検討

不適切な理由と、正解になり得る条件を対にして確認します。

A 不正解

ALBはHTTP/HTTPSなどL7向けで、任意TCPとUDPリスナーを提供しない。ノードIPも固定許可リスト用として管理する設計ではない。

この選択肢が正解になる条件

HTTP/HTTPS APIでホスト・パス・ヘッダー条件ルーティング、WAF、HTTP認証連携が必要なならALBが正解になり得る。

B 不正解

CLBへElastic IPを割り当てる構成やUDPを同一リスナーへカプセル化する説明は成立せず、新規設計でNLBより要件適合性が低い。

この選択肢が正解になる条件

既存のEC2-Classic世代アプリケーション互換を維持するなど、移行制約が強くTCP/HTTPの基本負荷分散だけならCLB継続が候補になり得る。

C 不正解

GWLBはGENEVEを用いてファイアウォールなど仮想アプライアンスを透過的にスケールするサービスで、一般クライアントへTCP/UDPアプリを直接公開する入口ではない。

この選択肢が正解になる条件

多数VPCのトラフィック経路へサードパーティ製ファイアウォールやIDSを透過挿入し、アプライアンスをスケールすることが目的なら正解になり得る。

D 正解

L4のTCP/UDP、高性能、静的IP、Multi-AZ、送信元IP保持をすべて満たす。

7 今後使える判断基準

ロードバランサーは、HTTPの内容で振り分けるならALB、TCP/UDP・静的IP・低遅延ならNLB、セキュリティアプライアンスの透過挿入ならGWLBと、プロトコルと役割から選ぶ。

8 関連サービスとの比較

ALBはHTTP/HTTPS/gRPC/WebSocketのL7機能、NLBはTCP/UDP/TLSのL4性能と静的IP、GWLBはGENEVEで仮想ネットワークアプライアンスを分散する。Global Acceleratorも静的Anycast IPを提供するが、リージョン内ターゲット分散はNLB等と組み合わせる。

QUESTION 016

問題16

COST

本番相当

複数選択 (2つ)

不確実・確実なアクセスパターンのS3階層化

企業は2種類のS3データを持つ。データXは平均20MBの設計ファイル100TBで、案件により数か月アクセスされないことも、突然毎日使われることもある。アクセス時はミリ秒で取得したく、取出し料金の予測不能を避けたい。データYは平均5MBの監査ログで、作成後30日間は調査に使うが、それ以降は7年間ほぼ読まれず、監査時の復元に48時間を許容する。法令により7年より前の削除は禁止される。直近2年を調べてもXに共通の年齢ベース規則は見つからず、Yの利用実績は日数条件と一致している。両データは数百万個以上のオブジェクトから成るため、監視料金と取出し料金も無視できない。運用担当者はオブジェクトごとのアクセス日を分析せず、既知の時系列ルールだけは自動化したい。保存費を最小化するために採用すべき方針を2つ選択してください。

- A** データXを作成直後からS3 Standard-IAへ保存し、アクセス頻度が上がっても同じクラスに固定する。最小保存期間と取出し料金はアクセス予測不能性より優先しない。
- B** データXをS3 Intelligent-Tieringへ移行し、アクセスパターンに応じて低遅延アクセス階層間を自動移動させる。オブジェクト監視料金を含めて評価する。
- C** データYのバケットでVersioningとS3 Object LockのComplianceモードを有効にし、各版を7年間保持する。Lifecycleで30日後にS3 Glacier Deep Archiveへ移行し、保持期間満了後に期限切れ削除する。
- D** データYをS3 Intelligent-TieringのFrequent Access階層に固定し、自動アーカイブ階層を無効にする。7年間は常にミリ秒アクセスを維持する。

ここで答えを決める

正解・解説は次のページから

ANSWER 016

問題16の正解・解説

1 正解 B・C

- B** データXをS3 Intelligent-Tieringへ移行し、アクセスパターンに応じて低遅延アクセス階層間を自動移動させる。オブジェクト監視料金を含めて評価する。
- C** データYのバケットでVersioningとS3 Object LockのComplianceモードを有効にし、各版を7年間保持する。Lifecycleで30日後にS3 Glacier Deep Archiveへ移行し、保持期間満了後に期限切れ削除する。

2 問題文の重要な条件

- Xはアクセス時期が予測不能だがミリ秒取得が必要
- Yは30日後ほぼ不使用、7年保持、復元48時間許容という既知パターンである
- 手動分析を避け、取出し・監視・最低保存期間を含む総費用を下げる

3 正解へ至る判断過程

1. 予測不能なアクセスと、日数で明確に予測できるアクセスを同じ階層化方式にしない。
2. XはIntelligent-Tieringの自動階層化で低遅延を維持しながら利用変化へ追従し、YはLifecycleで最安の長期アーカイブへ決定的に移せる。
3. Yの48時間復元はDeep Archiveに合い、Object LockのCompliance保持で7年未満の削除を技術的に防げる。XをStandard-IA固定すると取出し費用と最低保存期間のリスクが残る。

4 正解が要件を満たす理由

Bはアクセス時期を予測できないXをオブジェクト単位で自動追従させ、低遅延階層間の移動では取出し性能を維持する。CはYをLifecycleでDeep Archiveへ移して長期保存単価を下げつつ、VersioningとObject LockのCompliance保持により管理者を含め7年未満の削除を防ぐ。保持満了後の期限切れ削除まで自動化できる。

5・6 各選択肢の検討

不適切な理由と、正解になり得る条件を対にして確認します。

A 不正解

Standard-IA固定は再び頻繁に使われるXでも取出し料金が発生し、短期間で移動・削除する場合は最低保存期間の影響もある。変動へ自動追従しない。

この選択肢が正解になる条件

各オブジェクトが少なくとも所定の最低期間保存され、月に数回以下しかアクセスされないことを予測でき、取出し料金も見積もれるなら正解になり得る。

B 正解

Xの予測不能な変化に管理作業なしで追従し、必要なミリ秒アクセスを維持できる。

C 正解

Yの既知の時系列と長い復元許容をLifecycleとDeep Archiveへ対応させ、Object LockのCompliance保持で早期削除も防ぐ。

D 不正解

Yは30日後ほぼ読まれず48時間復元を許容するため、7年間低遅延の高い階層へ固定するとアーカイブによる節約を失う。

この選択肢が正解になる条件

監査ログを予告なしにミリ秒で読み出す法的要件があり、アーカイブ復元待ちを許容できないなら低遅延階層の維持が正解になり得る。

7 今後使える判断基準

アクセス時期が変動・予測不能ならIntelligent-Tiering、作成後N日で利用価値が明確に変わるならLifecycleを使う。ただし『保存する』と『削除を禁止する』は別要件であり、後者にはVersioningとObject Lockの保持モードを組み合わせる。

8 関連サービスとの比較

S3 Intelligent-Tieringは監視料金と引き換えにアクセス変化へ自動追従し、Lifecycleは年齢等で遷移・削除を実行する。Glacier Deep Archiveは最安級だが非同期復元を伴う。Object Lock Complianceは保持中の版をrootユーザーを含め削除・短縮できなくし、単なるLifecycle設定より強い不変性を提供する。

QUESTION 017

問題17

SECURE

本番より難しい

単一選択

高可用なRDS資格情報ローテーション

医療SaaSは、ECS上のアプリケーションからRDS for PostgreSQLへ接続している。現在のDBパスワードはタスク定義の環境変数へ手動登録され、半年以上変更されていない。監査対応として30日ごとに自動ローテーションし、平文の長期資格情報をイメージやIaCへ含めない構成へ変更したい。アプリケーションは高頻度に新規接続を作り、資格情報更新時にも新規接続の失敗を最小限にする必要がある。接続プールの古い接続は完了まで使えてよいが、更新直後の新規タスクも有効な値を取得する必要がある。監査ログには誰がシークレットを取得・変更したかを残したい。DBユーザー権限は必要最小限とし、アプリケーションは再デプロイせず最新資格情報を取得できなければならない。運用担当者によるDBと保存値の二重更新も避けたい。最も適切な構成はどれか。

- A** アプリ用の最小権限DB資格情報をAWS Secrets Managerへ保存し、30日周期の `alternating users rotation` を設定する。ECSタスクロールに取得権限を与え、アプリは接続作成時に現在のシークレットを取得・キャッシュ更新する。
- B** パスワードをSystems Manager Parameter StoreのStandard SecureStringへ保存し、EventBridge Schedulerから30日ごとに値だけをランダム更新する。RDS側のパスワードは半年ごとに手動で合わせる。
- C** RDSマスターパスワードをCloudFormationのNoEchoパラメータに置き、30日ごとにスタックを更新して全ECSタスクを再デプロイする。古いタスクは直ちに停止する。
- D** 暗号化済みパスワードをコンテナイメージへ埋め込み、復号用KMSキーをECSタスクロールへ許可する。30日ごとに新しいイメージを作成して全環境へローリングデプロイする。

ここで答えを決める

正解・解説は次のページから

ANSWER 017

問題17の正解・解説

1 正解 **A**

- A** アプリ用の最小権限DB資格情報をAWS Secrets Managerへ保存し、30日周期のalternating users rotationを設定する。ECSタスクロールに取得権限を与え、アプリは接続作成時に現在のシークレットを取得・キャッシュ更新する。

2 問題文の重要な条件

- RDS側と保存側を30日ごとに自動で同期ローテーションする
- 高頻度の新規接続でローテーション中の失敗を最小化する
- 最小権限・平文非埋込み・アプリ再デプロイ不要を満たす

3 正解へ至る判断過程

1. 資格情報の安全な保管だけでなく、DB内のパスワードとシークレット値を一体で更新する機能が必要である。
2. Secrets ManagerのDBローテーションを使い、高可用性が重要なアプリでは2ユーザーを交互に更新する方式を選ぶ。
3. タスク実行ロールではなくアプリのタスクロールへGetSecretValueを最小付与し、アプリ側も永続キャッシュせず更新を再取得できるようにする。

4 正解が要件を満たす理由

AはSecrets ManagerがシークレットとRDSユーザーの資格情報を同じローテーション処理で更新するため、手動二重更新をなくす。alternating users方式では片方の資格情報を更新している間もう片方が有効で、新規接続の可用性を高められる。タスクロールと動的取得により、イメージ・IaCへの埋込みや再デプロイも不要になる。

5・6 各選択肢の検討

不適切な理由と、正解になり得る条件を対にして確認します。

A 正解

DB連携ローテーション、高可用方式、動的取得、最小権限をまとめて満たす。

B 不正解

Parameter Storeは安全な値保管には使えるが、値だけを更新してRDS側を同時更新しないため認証が不一致になる。半年ごとの手動同期も30日自動化要件に反する。

この選択肢が正解になる条件

対象がDB資格情報ではない静的設定値で、外部システム側のローテーションが不要、変更通知と取得だけを低コストで管理したいなら正解になり得る。

C 不正解

NoEchoは表示抑制でシークレット管理・自動DBローテーションの代替ではない。マスター資格情報のアプリ利用は過剰権限で、30日ごとの全再デプロイと即時停止は接続失敗を増やす。

この選択肢が正解になる条件

初期構築時だけ機密パラメータを受け渡し、実行時資格情報は別のSecrets Manager参照へ移すならNoEchoが補助策になり得る。

D 不正解

暗号化されていても長期資格情報をイメージへ同梱し、更新のたびにビルド・配布するため、漏えい範囲と運用負荷が大きい。RDS側更新も記述されていない。

この選択肢が正解になる条件

資格情報ではなく署名済みの不変アセットをイメージへ含め、リリースと同時にのみ更新する要件ならKMS暗号化物の同梱が候補になり得る。

7 今後使える判断基準

DB資格情報では『保存』と『DB側の更新』を一体で考え、定期自動ローテーションならSecrets Managerを選ぶ。接続可用性が高いアプリはalternating users、単純性重視や対話ユーザーはsingle userを比較する。

8 関連サービスとの比較

Secrets Managerは取得・監査に加えてDB資格情報の自動ローテーションを提供する。Parameter Store SecureStringは設定・機密値保管に適するが、DBユーザー更新を自動で連携する標準ワークフローではない。IAM DB authenticationは対応エンジンでパスワード代替も検討できる。

QUESTION 018

問題18

RESILIENT

複合難問

単一選択

SQSバーストから下流DBを保護

受注システムはSQS標準キューをLambdaで処理し、各実行がAuroraへ最大2本の接続を開く。同じLambda関数は同期APIからも呼び出され、その経路には最大30同時実行を確保したい。セール時にキューへ50万件が到着すると、SQSイベントソースマッピングが急速にスケールし、DBの安全な接続上限100を超えてタイムアウトが連鎖する。メッセージ処理は数分遅れてもよく、キューに耐久的に滞留させられる。1バッチの最長処理時間は40秒で、失敗メッセージは他の注文を永続的に妨げないよう隔離したい。既存関数のコード分割は今回のリリースではできない。アカウント全体のLambda同時実行を不必要に下げず、同期APIへの容量も奪わず、SQS経路だけを最大20同時実行に抑えたい。設定変更だけで行う。最も適切な構成はどれか。

- A** 関数のreserved concurrencyを20に設定し、SQSイベントソースのmaximum concurrencyは未設定にする。同期APIも同じ20枠を共有し、キューの可視性タイムアウトは変更しない。
- B** 関数のreserved concurrencyを少なくとも50に設定し、SQSイベントソースマッピングのmaximum concurrencyを20にする。処理時間に合う可視性タイムアウトとDLQを設定し、DB接続数も監視する。
- C** SQSキューの配信遅延を最大値にし、Lambdaの同時実行は無制限のままにする。バースト時は全メッセージが同じ時間だけ遅れるためDB接続数も常に20以下になると見込む。
- D** Auroraの最大接続数を直ちに10倍へ変更し、SQSイベントソースマッピングをProvisioned Modeの最大ポーラー数で動かす。Lambda側の同時実行制御は行わない。

ここで答えを決める

正解・解説は次のページから

ANSWER 018

問題18の正解・解説

1 正解 **B**

- B** 関数のreserved concurrencyを少なくとも50に設定し、SQSイベントソースマッピングのmaximum concurrencyを20にする。処理時間に合う可視性タイムアウトとDLQを設定し、DB接続数も監視する。

2 問題文の重要な条件

- SQS経路だけを最大20同時実行に制限する
- 同じ関数の同期API用に最大30同時実行を残す
- メッセージ遅延は許容し、SQSの耐久バッファと再試行を活用する

3 正解へ至る判断過程

1. 関数全体のreserved concurrencyと、特定イベントソースのmaximum concurrencyの適用範囲を区別する。
2. SQSマッピングだけを20へ制限し、関数全体にはSQS 20と同期API 30を収容できる上限を確保すれば、別経路を阻害しない。
3. 処理が遅れる分はSQSが吸収し、可視性タイムアウト・DLQ・監視で重複処理や毒メッセージによる枠占有を抑える。

4 正解が要件を満たす理由

Bはイベントソースマッピング単位のmaximum concurrencyでSQSポーリングからの同時Lambda実行を20に抑え、最大40本程度のDB接続へ制限できる。関数全体のreserved concurrencyを50以上にすることで同期API用30枠との合計も収容する。キューがバーストを保持するためデータを落とさず、DB能力に合わせて排出できる。

5・6 各選択肢の検討

不適切な理由と、正解になり得る条件を対にして確認します。

A 不正解

DB保護自体はできるが、関数全体を20に制限するためSQSが枠を使うと同期APIの最大30を確保できない。可視性タイムアウト不足は処理中メッセージの再出現も招く。

この選択肢が正解になる条件

関数がSQS専用で、他の呼出経路がなく、全経路合計を20へ制限したいならreserved concurrency 20が正解になり得る。

B 正解

SQS経路だけを絞りながら関数全体にはAPI容量を残し、キューの耐久性で安全にバックプレッシャーを作る。

C 不正解

配信遅延は各メッセージの可視化時刻をずらす、同時実行レートを20へ制御する仕組みではない。大量メッセージが同時に可視化されれば再び急増する。

この選択肢が正解になる条件

個々の処理を一定時間後に開始するスケジュール要件だけで、下流容量制御が不要なら配信遅延が正解になり得る。

D 不正解

DB接続上限の拡大はメモリ・CPUとの関係を見落とし、Provisioned Modeは取り込みを速める方向で障害を悪化させ得る。下流能力に合わせるバックプレッシャーがない。

この選択肢が正解になる条件

DBが実測上より多い接続と処理量を安全に扱え、目標がSQS処理開始の急速なスケールアップならProvisioned Modeが正解になり得る。

7 今後使える判断基準

Lambdaの下流保護では、関数全体を絞るreserved concurrencyと、SQSマッピング単位を絞るmaximum concurrencyを使い分ける。キューは遅延を吸収し、可視性タイムアウトとDLQを併設する。

8 関連サービスとの比較

reserved concurrencyは関数の全呼出しに対する保証兼上限、SQS maximum concurrencyはそのイベントソースからの上限である。SQSはバックプレッシャーを耐久保持し、RDS Proxyは接続再利用に有効だが、DBのクエリ処理能力そのものを無限に増やすものではない。

QUESTION 019

問題19

PERFORMANCE

本番相当

単一選択

動的TCP・UDPのグローバル高速化

ゲーム会社は東京とフランクフルトのNLBで、リアルタイム対戦サーバーを提供している。クライアントはUDPの位置更新とTCPの制御接続を使用し、コンテンツは利用者ごとに变化するためキャッシュできない。南米と欧州の利用者から、パブリックインターネット区間のジッターとパケット損失が報告されている。クライアントソフトには少数の固定IPv4アドレスを埋め込みたい。いずれかのリージョンやNLBが不健全なら、正常なリージョンへ新規接続を自動誘導したい。既存NLBとポートは維持し、クライアントから最初のAWS入口までのインターネット距離を短縮することも求められる。静的アセット配信は別のCDNで既に解決済みである。DNS TTLに依存した切替やHTTPへのプロトコル変更は避ける。接続先は自動判定する。最も適切な構成はどれか。

- A** 各NLBの前にCloudFrontディストリビューションを配置し、UDPと任意TCPをキャッシュ無効で転送する。CloudFrontのエッジIPをクライアントへ固定登録する。
- B** Route 53のレイテンシールーティングで2つのNLBを返し、TTLを1秒にする。クライアントには最初に解決したNLBの現在のIPを端末内へ永続保存させる。
- C** AWS Global Acceleratorの標準アクセラレーターを作成し、TCP/UDPリスナーと両リージョンのNLBエンドポイントを設定する。提供される静的Anycast IPをクライアントに登録し、ヘルスに基づき経路選択する。
- D** S3 Transfer Accelerationを有効にし、ゲームパケットを各リージョンのS3バケットへアップロードする。Lambdaがオブジェクト作成イベントから対戦サーバーへ転送する。

ここで答えを決める

正解・解説は次のページから

ANSWER 019

問題19の正解・解説

1 正解 C

- C** AWS Global Acceleratorの標準アクセラレーターを作成し、TCP/UDPリスナーと両リージョンのNLBエンドポイントを設定する。提供される静的Anycast IPをクライアントに登録し、ヘルスに基づき経路選択する。

2 問題文の重要な条件

- キャッシュ不能な動的UDPとTCPを高速化する
- 固定IPを提供し、DNS TTL依存なしで正常リージョンへ誘導する
- インターネット区間を短くし、AWSグローバルネットワークを活用する

3 正解へ至る判断過程

1. CloudFrontはHTTP系コンテンツ配信であり、任意UDPゲーム通信のプロキシではない。
2. Global Acceleratorは静的Anycast IPへ最寄りエッジから入り、TCP/UDPをAWSネットワーク経由で正常なリージョナルエンドポイントへ送る。
3. DNSベースのRoute 53より固定IPと迅速なヘルスルーティング条件に合い、キャッシュ可能性にも依存しない。

4 正解が要件を満たす理由

CはTCPとUDPの動的通信を対象に、クライアントに変更されないAnycast IPを提供する。利用者から最寄りのAWSエッジへ取り込み、AWSグローバルネットワーク上で正常なNLBへ送るため、インターネット区間のジッターを減らせる。エンドポイントヘルスに基づく切替はDNSキャッシュ更新を待たない。

5・6 各選択肢の検討

不適切な理由と、正解になり得る条件を対にして確認します。

A 不正解

CloudFrontはHTTP/HTTPS配信とHTTP上のWebSocketなどに向くが、任意UDPと独自TCPリスナーをそのまま転送しない。エッジIP一覧もクライアント固定用の専用静的IPとは異なる。

この選択肢が正解になる条件

静的・動的なHTTP/HTTPSコンテンツ、API、動画セグメントをエッジキャッシュや接続最適化で配信するならCloudFrontが正解になり得る。

B 不正解

Route 53は正常・低遅延リージョンのDNS選択には使えるが、パケットをAWSバックボーンへ取り込まず、TTLを下げてクライアントや中間リゾルバーのキャッシュに依存する。IP永続保存は切替を妨げる。

この選択肢が正解になる条件

DNS名利用を許容し、リージョン選択だけが必要で、ネットワーク高速化や固定IPが不要ならレイテンシールーティングが正解になり得る。

C 正解

TCP/UDP、静的Anycast IP、AWSネットワーク経由の高速化、正常エンドポイントへの切替を一体で満たす。

D 不正解

S3 Transfer Accelerationはオブジェクト転送をエッジ経由で高速化する機能で、対話的なゲームパケットや長時間TCP接続の中継ではない。オブジェクトイベント化は遅延を大幅に増やす。

この選択肢が正解になる条件

世界中の利用者が大きなS3オブジェクトを単一バケットへアップロード・ダウンロードし、その転送時間短縮が目的なら正解になり得る。

7 今後使える判断基準

キャッシュ可能なHTTP配信ならCloudFront、動的なTCP/UDPを静的IPでグローバル高速化するならGlobal Accelerator、DNSでリージョンだけ選ぶならRoute 53と切り分ける。

8 関連サービスとの比較

Global AcceleratorはL4のTCP/UDP、静的Anycast IP、正常エンドポイントへの経路最適化を提供する。CloudFrontはL7のHTTP/HTTPS配信とキャッシュ、Route 53はDNS回答制御、S3 Transfer AccelerationはS3オブジェクト転送専用である。

QUESTION 020

問題20

COST

本番より難しい

複数選択 (2つ)

復元時間と保持期間によるGlacier選択

保険会社は2種類のS3アーカイブを分類している。データXは災害復旧用バックアップで、1TBのオブジェクト群を最低1年間保持し、年1〜2回の訓練で復元する。通常は3〜5時間以内に利用開始できればよいが、重大障害時には追加料金を許容して一部を数分で取り出せる選択肢も求める。データYは契約終了済み記録で、法令により10年間保存し、通常は一度も読まれない。監査時の復元は48時間以内でよく、とにかく保存単価を下げたい。年間保存量はYがXの20倍で、取出し回数より保存容量が請求を支配している。復元後の一時コピーを使うアプリ改修は完了している。どちらもミリ秒アクセスは不要で、オブジェクトを早期削除する予定はなく、各クラスの最低保存期間より十分長く保持する。最適な保存クラスを2つ選択してください。

- A** データXにはS3 Glacier Instant Retrievalを選び、常にミリ秒で直接取得できる状態を維持する。年1〜2回のアクセスでも復元ジョブは使用しない。
- B** データXにはS3 Glacier Flexible Retrievalを選ぶ。通常はStandardまたはBulk、重大障害時の一部データには利用可能なExpedited retrievalを使い分ける。
- C** データYにはS3 Glacier Flexible Retrievalを選び、10年間すべて同じクラスに置く。3〜5時間のStandard retrievalを常に確保するため保存単価差は考慮しない。
- D** データYにはS3 Glacier Deep Archiveを選ぶ。StandardまたはBulkの復元時間を48時間要件内で使い分け、180日より十分長い10年保持に適用する。

ここで答えを決める

正解・解説は次のページから

ANSWER 020

問題20の正解・解説

1 正解 **B・D**

- B** データXにはS3 Glacier Flexible Retrievalを選ぶ。通常はStandardまたはBulk、重大障害時の一部データには利用可能なExpedited retrievalを使い分ける。
- D** データYにはS3 Glacier Deep Archiveを選ぶ。StandardまたはBulkの復元時間を48時間要件内で使い分け、180日より十分長い10年保持に適用する。

2 問題文の重要な条件

- Xは年1～2回、通常3～5時間、緊急時は一部を数分で復元したい
- Yは10年ほぼ無アクセスで、48時間以内の復元と最低保存単価を求める
- どちらも最低保存期間を大幅に超えて保持し、ミリ秒アクセスは不要

3 正解へ至る判断過程

1. Glacier系は保存単価だけでなく、直接アクセスか復元ジョブか、利用できる復元ティア、最低保存期間で選ぶ。
2. XはFlexible RetrievalのStandard 3～5時間が通常要件に合い、必要な一部ではExpeditedという選択肢を持てる。
3. YはDeep ArchiveのStandard/Bulkが48時間以内に収まり、10年保存なので180日の最低保存期間を問題なく超え、最安の長期保管を活かせる。

4 正解が要件を満たす理由

BはXの通常復元時間と緊急時の数分復元を、Flexible Retrievalの複数復元ティアで費用と速度を使い分ける。Dはほぼ読まれないYを最も低コストのDeep Archiveへ置き、最大48時間という緩い復元要件と10年保持を活用する。各データを必要以上に高速なクラスへ置かない。

5・6 各選択肢の検討

不適切な理由と、正解になり得る条件を対にして確認します。

A 不正解

Instant Retrievalはミリ秒の直接アクセスが必要なアーカイブ向けだが、Xは通常3～5時間を許容し年1～2回だけなので、常時低遅延に支払う必要がない。

この選択肢が正解になる条件

医用画像など、アクセス頻度は低いが要求時には復元ジョブを待たずミリ秒で取得しなければならないなら正解になり得る。

B 正解

通常と緊急で復元ティアを選べ、Xの頻度・時間・1年保持に最も合う。

C 不正解

Flexible RetrievalでもYの48時間要件は満たすが、10年ほぼ無アクセスで保存単価最優先ならDeep Archiveの方が適する。不要な3～5時間性能へ費用を払う。

この選択肢が正解になる条件

Yを年1～2回読み、復元を通常3～5時間、場合によって数分に短縮する必要があるなら正解になり得る。

D 正解

Yの非常に長い保持、ほぼ無アクセス、48時間復元という条件に最安級の長期アーカイブを対応させる。

7 今後使える判断基準

Glacier選択では、ミリ秒直接取得ならInstant Retrieval、数分～数時間を使い分けるならFlexible Retrieval、12～48時間待てる長期ほぼ無アクセスならDeep Archiveを基準にし、90日・180日などの最低保存期間も確認する。

8 関連サービスとの比較

Glacier Instant Retrievalは低頻度だがミリ秒アクセス、Flexible Retrievalは復元ジョブを使い数分～時間、Deep Archiveは最安級で通常12時間・Bulk最大48時間程度を想定する。Standard-IAはアーカイブ復元なしの即時アクセスだが長期最低コスト用途とは異なる。

QUESTION 021

問題21

SECURE

複合難問

単一選択

双方向ハイブリッドDNS

製造会社はDirect Connectでオンプレミスと複数のAWSアカウントのVPCを接続している。オンプレミスDNSはcorp.example.comを権威管理し、VPC内のRoute 53 private hosted zoneはaws.corp.example.comの名前を保持する。EC2からdb.corp.example.comを解決し、オンプレミス端末からapi.aws.corp.example.comも解決できなければならない。DNSクエリはインターネットへ出さず、既存の中央ネットワークアカウントで集約し、ルールをAWS RAMで各アカウントへ共有したい。単一AZのDNS中継障害で名前解決が止まることも許されない。自前のBINDサーバー、条件転送EC2、手動フェイルオーバーは運用したくない。最も適切な構成はどれか。

- A** 各VPCのDHCP options setにオンプレミスDNSだけを指定し、private hosted zoneのレコードもオンプレミスDNSへ毎時間ゾーン転送する。Route 53 Resolverは使用しない。
- B** 中央VPCにRoute 53 Resolver inbound endpointだけを複数AZで作成する。EC2からオンプレミスドメインへのクエリも同じinbound endpointへ転送する。
- C** 中央VPCにRoute 53 Resolver outbound endpointだけを作成し、corp.example.comの転送ルールを各VPCへ共有する。オンプレミスからprivate hosted zoneへの問い合わせは公開hosted zoneへ複製する。
- D** 中央VPCの複数AZにResolver inbound endpointとoutbound endpointを作る。corp.example.comはoutbound ruleでオンプレミスDNSへ送り、オンプレミスDNSはaws.corp.example.comをinbound endpointへ条件転送し、ルールをRAM共有する。

ここで答えを決める

正解・解説は次のページから

ANSWER 021

問題21の正解・解説

1 正解 **D**

- D** 中央VPCの複数AZにResolver inbound endpointとoutbound endpointを作る。
corp.example.comはoutbound ruleでオンプレミスDNSへ送り、オンプレミスDNSは
aws.corp.example.comをinbound endpointへ条件転送し、ルールをRAM共有する。

2 問題文の重要な条件

- AWSからオンプレミス、オンプレミスからprivate hosted zoneの双方向名前解決が必要
- Direct Connect上の私設経路を使い、複数アカウントへ中央ルールを共有する
- 自前DNS転送サーバーを持たず、エンドポイントを複数AZ化する

3 正解へ至る判断過程

1. 問い合わせ方向ごとにResolver endpointの役割を分ける。VPCから外へ出すのがoutbound、オンプレミスからVPC Resolverへ入れるのがinboundである。
2. オンプレミス権威ドメイン用のforward ruleをoutbound endpointへ関連付け、private hosted zone用にはオンプレミスDNSからinbound endpoint IPへ条件転送する。
3. 中央VPCで複数AZに配置し、Resolver ruleをRAM共有すれば、高可用性とマルチアカウント運用を両立できる。

4 正解が要件を満たす理由

Dは双方向をそれぞれ正しいエンドポイントで処理する。EC2のcorp.example.com問い合わせは共有されたoutbound ruleからDirect Connect経由でオンプレミスDNSへ行き、オンプレミスのaws.corp.example.com問い合わせはinbound endpointからVPC Resolverとprivate hosted zoneへ届く。複数AZ配置とRAM共有で、独自DNSサーバーなしに可用性・集中管理も満たす。

5・6 各選択肢の検討

不適切な理由と、正解になり得る条件を対にして確認します。

A 不正解

VPCのAmazonProvidedDNSを外す設計はAWS内部名やprivate hosted zoneの解決を複雑にし、Route 53 private hosted zoneの一般的なゾーン転送を前提にしている点も不適切である。

この選択肢が正解になる条件

全名前空間をオンプレミスDNSだけが権威管理し、AWS固有のprivate hosted zoneを使わず、既存DNS冗長化を運用する明確な要件ならカスタムDNS指定が候補になり得る。

B 不正解

inbound endpointはオンプレミスからVPC Resolverへの入口であり、VPC発のクエリをオンプレミスへ転送する用途ではない。outbound側が欠ける。

この選択肢が正解になる条件

必要なのがオンプレミスからprivate hosted zoneを引く一方向だけで、AWSからオンプレミス名を解決しないならinbound endpointだけで正解になり得る。

C 不正解

outbound側はAWSからオンプレミス名を解決できるが、オンプレミスからprivate hosted zoneへ入る経路がない。非公開名を公開hosted zoneへ複製すると情報公開と二重管理が生じる。

この選択肢が正解になる条件

必要なのがVPCからオンプレミス名を引く一方向だけで、オンプレミスからAWS非公開名を解決しないなら正解になり得る。

D 正解

inbound/outboundを双方向に使い分け、私設接続、Multi-AZ、ルール共有、管理サービス利用をすべて満たす。

7 今後使える判断基準

ハイブリッドDNSでは問い合わせの矢印を書く。VPC→オンプレミスはoutbound endpointとforward rule、オンプレミス→private hosted zoneはinbound endpointとオンプレミス側条件転送である。

8 関連サービスとの比較

Resolver inbound endpointは外部DNSからVPC Resolverへの入口、outbound endpointはVPC Resolverから指定DNSへの出口である。Private hosted zoneは関連VPC内の非公開名前空間、public hosted zoneはインターネット公開用で、非公開レコードの代替公開は避ける。

QUESTION 022

問題22

RESILIENT

本番相当

単一選択

Route 53によるactive-passive切替

予約サービスは東京リージョンのALBをプライマリ、大阪リージョンのALBをウォームスタンバイとしている。データは別途複製され、両リージョンのアプリケーションは/health/deepでDB接続を含む正常性を返す。通常は全利用者を東京へ送り、東京のALBが応答していてもアプリ依存関係が連続して失敗した場合は大阪へDNSフェイルオーバーしたい。大阪が不健全なときに誤って切り替えず、東京復旧後は自動的にプライマリへ戻したい。両ALBのノードIPは変化し得るため、IP一覧を定期同期する運用は採用できない。切替はDNSキャッシュの範囲内によく、既存の単一ドメイン名を維持する。クライアント変更や独自の監視Lambda、手動DNS更新は避ける。監視もAWSに任せたい。最も適切なRoute 53構成はどれか。

- A** 同じ名前にPrimaryとSecondaryのfailover alias recordを作成し、両recordでEvaluate Target Healthを有効にする。各ALBのターゲットグループは/health/deepをヘルスチェックし、正常なPrimaryを優先して失敗時だけ正常なSecondaryを返す。
- B** 2つのALBに対するweighted recordを各50に設定し、ヘルスチェックは作成しない。東京障害時は大阪の重みを運用担当者が手動で100へ変更する。
- C** latency-based routing recordを作成し、Evaluate Target Healthを無効にする。通常時も利用者に近いリージョンへ送り、リージョン障害はクライアント再試行で処理する。
- D** multivalue answer recordで両ALBのIPアドレスを直接登録し、最大8件を常に返す。ALBのIP変更はEventBridgeで検出してレコードへ反映する。

ここで答えを決める

正解・解説は次のページから

ANSWER 022

問題22の正解・解説

1 正解 **A**

- A** 同じ名前にPrimaryとSecondaryのfailover alias recordを作成し、両recordでEvaluate Target Healthを有効にする。各ALBのターゲットグループは/health/deepをヘルスチェックし、正常なPrimaryを優先して失敗時だけ正常なSecondaryを返す。

2 問題文の重要な条件

- 通常は東京だけを使うactive-passiveである
- ALBだけでなくDB接続を含む深い正常性で自動切替する
- Secondaryの正常性も確認し、自動フェイルバックと低運用負荷を求める

3 正解へ至る判断過程

1. 平常時の優先順位が固定されたactive-passiveなので、weightedやlatencyではなく failover routingを選ぶ。
2. 各ALBのターゲットグループでアプリ依存関係を含むURLを確認し、failover alias recordの Evaluate Target Healthで両リージョンの正常性をDNS回答判断に使う。
3. AliasでALBを参照すれば変動するロードバランサーIPを管理せず、Primary回復後の優先回答も自動化できる。

4 正解が要件を満たす理由

AはPrimary/Secondaryという明示的な優先順位を持ち、ALBターゲットグループの深いヘルスチェック結果をEvaluate Target HealthでRoute 53へ反映する。Primaryが不健全な場合だけ正常なSecondaryを返し、東京が回復すれば自動的にPrimaryを再び優先する。ALBのalias recordによりIP追跡や独自Lambdaも不要である。

5・6 各選択肢の検討

不適切な理由と、正解になり得る条件を対にして確認します。

A 正解

active-passiveの優先順位、アプリレベル正常性、Secondary確認、自動フェイルバックをRoute 53の標準機能で満たす。

B 不正解

50/50 weightedは通常時から大阪へ送るためactive-passiveに反し、ヘルスチェックなし・手動重み変更では自動RTOも満たさない。

この選択肢が正解になる条件

両リージョンをactive-activeで利用し、カナリアや段階移行のため割合を制御し、各recordに正常性評価も設定するならweighted routingが正解になり得る。

C 不正解

latency routingは利用者ごとに低遅延リージョンを選ぶactive-active用途で、東京常用という要件に合わない。正常性評価を無効にすると障害リージョンも回答され得る。

この選択肢が正解になる条件

両リージョンが常時書込み可能で、利用者を最小遅延の正常リージョンへ送るactive-active設計なら正解になり得る。

D 不正解

ALBの変動IPを直接登録・追跡する必要はなく、multivalueはPrimary/Secondaryの優先フェイルオーバーを表さない。独自更新運用も制約に反する。

この選択肢が正解になる条件

複数の独立した固定IP Webサーバーから正常な最大8件をランダムに返す簡易DNS負荷分散が目的なら正解になり得る。

7 今後使える判断基準

Route 53では通常時の配送モデルから選ぶ。active-passiveはfailover、比率配分はweighted、最小遅延のactive-activeはlatency。DNS切替にはPrimaryだけでなくSecondaryの正常性とアプリ深度も含める。

8 関連サービスとの比較

Failover routingはPrimary優先のactive-passive、weightedは割合制御、latencyは低遅延リージョン選択、multivalue answerは複数の正常レコード返却に向く。ALBはalias targetにしてIPアドレスの変化をRoute 53へ任せる。

QUESTION 023

問題23

PERFORMANCE

本番より難しい

単一選択

ストリームの即時処理と再読込

決済会社は、毎秒2万件まで増える取引イベントをAWSへ取り込み、不正検知モデルと顧客向け通知で利用している。繁忙時には現在の配信基盤がイベントをまとめて送るため、検知開始まで30～90秒かかっている。新構成では、各イベントを通常2秒以内に独立した2つのコンシューマーへ渡す必要がある。一方のコンシューマーが6時間停止しても、復旧後に停止位置から再処理できなければならない。全イベントは監査用にAmazon S3へも保存する。初期のシャード数見積りは避けたいが、稼働前の負荷試験では同じストリームを段階的に毎秒2万件まで上げ、プロデューサーには指数バックオフ付き再試行を実装できる。イベントはパーティションキー単位の時系列を保ち、障害中もプロデューサーを止められない。S3配送コードの独自保守も減らしたい。運用負荷を抑えつつ、これらの要件を最もよく満たす構成はどれか。

- A** プロデューサーからAmazon Data Firehoseへ直接送信し、動的パーティショニングでS3へ配信する。2つのLambda変換をFirehoseの配信前処理として順番に実行し、失敗レコードはS3へ保存する。
- B** Kinesis Data Streamsをオンデマンドモード、保持期間72時間で作成する。2つの登録済みコンシューマーを拡張ファンアウトで接続し、別途Data FirehoseをコンシューマーとしてS3へ配信する。
- C** Amazon SNSの標準トピックから2つのSQS標準キューへファンアウトし、各キューの保持期間を72時間にする。処理済みイベントは各コンシューマーが個別に同じS3オブジェクトへ追記する。
- D** Amazon EventBridgeのカスタムイベントバスに全イベントを送信し、2つのLambdaターゲットとS3ターゲットを設定する。失敗時は再試行ポリシーだけで6時間分を再生する。

ここで答えを決める

正解・解説は次のページから

ANSWER 023

問題23の正解・解説

1 正解 **B**

- B** Kinesis Data Streamsをオンデマンドモード、保持期間72時間で作成する。2つの登録済みコンシューマーを拡張ファンアウトで接続し、別途Data FirehoseをコンシューマーとしてS3へ配信する。

2 問題文の重要な条件

- 同じストリームを段階的に毎秒2万件まで事前負荷試験し、シャード数の継続管理は避けたい
- 2つの独立コンシューマーへ通常2秒以内に渡し、読み取り競合を避ける
- 6時間停止後の位置から再処理し、同じデータをS3にも永続化する

3 正解へ至る判断過程

1. 単なる最終配送ではなく、複数コンシューマーが独立して低遅延に読み、一定期間後に再読込できるストリームが必要である。
2. Kinesis Data Streamsは保持期間内のレコードを再読込でき、拡張ファンアウトは登録済みコンシューマーごとに専用の読み取りスループットを与える。
3. オンデマンドモードでシャード運用を減らし、事前に段階的なピークを記録して再試行も備える。Data Firehoseを追加コンシューマーとして使えばS3配送の実装も抑えられる。

4 正解が要件を満たす理由

Bは低遅延のストリーム処理、独立した複数コンシューマー、72時間内のチェックポイントからの再処理、S3への管理された配送を同時に満たす。オンデマンドモードはシャード計画を減らし、拡張ファンアウトは読み取り帯域競合を避ける。さらに同じストリームを段階的に毎秒2万件まで到達させ、プロデューサー再試行を備えるため、過去ピークの2倍を大きく超える瞬時増加でスロットリングし得る点も織り込んでいる。

5・6 各選択肢の検討

不適切な理由と、正解になり得る条件を対にして確認します。

A 不正解

Data Firehoseは宛先への管理された配信には適するが、2つの独立コンシューマーが停止位置を管理して低遅延に再読込するログという要件の中心を満たさない。変換処理を直列化しても独立した購読にはならない。

この選択肢が正解になる条件

目的が数十秒程度のバッファを許容するS3への一方向配送と1種類の軽量変換だけで、任意位置からの再生や独立コンシューマーが不要なら正解になり得る。

B 正解

保持可能な低遅延ストリーム、専用読み取り帯域、容量自動管理、S3配送を組み合わせ、全条件を満たす。

C 不正解

SNSとSQSのファンアウトは処理分離には有効だが、ストリーム内のチェックポイントから複数コンシューマーが再生するモデルではない。また複数処理系から同一S3オブジェクトへ追記する設計は成立しない。

この選択肢が正解になる条件

各購読者を独立キューで疎結合化することが主目的で、厳密なストリーム再生や同一パーティション内の時系列処理が不要なら正解になり得る。

D 不正解

EventBridgeはイベントルーティングに適するが、再試行ポリシーは保持ストリームのオフセット管理の代替ではない。障害後に任意位置から6時間分を再処理する設計になっていない。

この選択肢が正解になる条件

イベント量がサービスクォータ内で、属性に基づく多数のSaaS・AWSターゲットへのルーティングが主目的かつ長時間の再生要件がなければ正解になり得る。

7 今後使える判断基準

複数の独立処理、秒単位の遅延、保持期間内の再生が必要ならKinesis Data Streamsを軸にする。ただしOn-Demandも直前ピークを大幅に超える急増を無条件には保証しないため、段階的な事前負荷、再試行・バッファ、クォータを確認する。S3への管理配送だけならData Firehoseを使う。

8 関連サービスとの比較

Kinesis Data Streamsはコンシューマーが保持レコードを読み再生するサービスで、拡張ファンアウトにより専用帯域を得られる。Data FirehoseはS3などの宛先への配送を管理する。SNS/SQSは購読者分離と耐久キュー、EventBridgeはイベント内容に基づくルーティングに強い。

QUESTION 024

問題24

COST

複合難問

複数選択 (2つ)

異なるI/O特性に対するEBS選択

映像解析会社は、EC2上の2種類のワークロードについてEBS費用を見直している。ワークロードXは、1台のインスタンスが毎晩8TBの一時ログを大きなブロックで順次読み書きし、ピーク時に約400MiB/秒を必要とする。ランダムI/Oは少なく、ブートボリュームには使わず、処理失敗時はS3の原本から再生成できる。ワークロードYは、単一EC2上の自己管理型取引DBで、ピーク時40,000 IOPS、ミリ秒未満の一貫した待ち時間、高い耐久性が必要である。容量はどちらも数TBで、性能不足は許されないが、不要なプロビジョンドIOPSへの支払いは避けたい。EC2側のEBS最適化帯域は必要性能を満たすインスタンスへ変更でき、バックアップ要件は別途満たしている。各ワークロードに最適なボリュームを選ぶ必要がある。最も適切な選択肢を2つ選択してください。

- A** ワークロードXにはsc1 Cold HDDを選び、必要な400MiB/秒は複数ボリュームのストライピングで補う。低頻度アクセスのため、スループット要件よりGB単価を優先する。
- B** ワークロードYにはgp3を選び、40,000 IOPSとミリ秒未満の一貫性をすべてgp3の追加IOPS設定だけで満たす。耐久性要件は日次スナップショットで補う。
- C** ワークロードXにはst1 Throughput Optimized HDDを選ぶ。大容量の順次I/O向けに必要な容量とスループットを見積もり、対応するEC2のEBS帯域も確認する。
- D** ワークロードYにはio2 Block Expressを選び、必要IOPSと容量をプロビジョニングする。対応インスタンスを使用し、DBのバックアップは別途取得する。

ここで答えを決める

正解・解説は次のページから

ANSWER 024

問題24の正解・解説

1 正解 C・D

- C** ワークロードXにはst1 Throughput Optimized HDDを選ぶ。大容量の順次I/O向けに必要な容量とスループットを見積もり、対応するEC2のEBS帯域も確認する。
- D** ワークロードYにはio2 Block Expressを選び、必要IOPSと容量をプロビジョニングする。対応インスタンスを使用し、DBのバックアップは別途取得する。

2 問題文の重要な条件

- Xは大容量・順次I/O・約400MiB/秒で、ブート用途ではない
- Yは40,000 IOPS、低く一貫した遅延、高耐久性が必要
- 性能を満たしつつ、不要なプロビジョンドIOPS費用を避ける

3 正解へ至る判断過程

1. EBSは容量だけでなく、I/Oが順次かランダムか、必要なIOPSかスループットか、耐久性の水準で分ける。
2. XはHDDに適した大きな順次I/Oで高スループットが必要なためst1、Yは高IOPSと一貫した低遅延を必要とするミッションクリティカルDBなのでio2 Block Expressが合う。
3. gp3は現行世代では40,000 IOPS自体を設定できるが、Yが求めるio2水準の高耐久性と一貫した低遅延まで日次スナップショットで代替できない。sc1は低頻度データ向けでXの継続的な高スループット優先と一致しない。

4 正解が要件を満たす理由

Cはst1の大容量・順次アクセス向け特性をXに対応させ、プロビジョンドIOPS SSDの過剰費用を避ける。DはYが求める40,000 IOPS、ミリ秒未満の一貫した遅延、高耐久性をio2 Block Expressで満たす。両方を選ぶことで、ワークロード特性ごとに性能と費用を最適化できる。

5・6 各選択肢の検討

不適切な理由と、正解になり得る条件を対にして確認します。

A 不正解

sc1は低頻度アクセス向けの低コストHDDであり、毎晩400MiB/秒を確実に必要とするXではst1の方が目的に合う。ストライピングは故障・監視・復旧の運用面も増やす。

この選択肢が正解になる条件

Xの処理が月数回以下で、要求スループットが大幅に低く、処理時間の延長を許容してGB単価を最優先できるならsc1が正解になり得る。

B 不正解

gp3は40,000 IOPSを設定できても、io2 Block Express向けの一貫した低遅延と高耐久性というYの要件を同等には満たさない。スナップショットは復旧手段であり、稼働中ボリュームの耐久性・遅延特性を置き換えない。

この選択肢が正解になる条件

40,000 IOPSと必要スループットがgp3の対応リージョン・構成上限内で、io2水準の一貫した低遅延や高耐久性が不要な一般業務DBなら、gp3が費用面で正解になり得る。

C 正解

順次アクセス中心の大容量スループット処理というXの特性に合い、SSDより費用を抑えられる。

D 正解

高IOPS、低く一貫した遅延、高耐久性を必要とするYに対応し、必要性能を明示的にプロビジョニングできる。

7 今後使える判断基準

EBS選定では、一般的なランダムI/Oはgp3、ミッションクリティカルな高IOPS・低遅延はio2、頻繁な大容量順次I/Oはst1、低頻度の大容量順次I/Oはsc1と切り分け、EC2側のEBS帯域も同時に確認する。

8 関連サービスとの比較

gp3は容量とIOPS・スループットを独立設定できる汎用SSD、io2 Block Expressは高IOPS・一貫した低遅延・高耐久性向けで高価である。st1は頻繁な順次スループット、sc1は低頻度アクセスの容量単価を重視するHDDで、HDD系はブートに使えない。

QUESTION 025

問題25

SECURE

本番相当

単一選択

クロスアカウントKMS復号の二層認可

企業はAWS Organizations配下で、データ所有アカウントAと分析アカウントBを分離している。AのS3バケット内のオブジェクトは、Aが所有するカスターマネージドKMSキーで暗号化されている。BのGlueジョブロールだけに復号を許可し、Bの他の管理者やロールには許可したくない。S3バケットポリシーによるオブジェクト読取許可はすでに正しく設定されているが、GlueジョブはKMSのAccessDeniedで失敗する。ジョブは一時的なSTS資格情報を使い、同じキーを別用途で利用する必要はない。CloudTrailでキー利用を追跡し、将来は権限を一箇所ずつ失効できるようにする。監査チームは最小権限を要求し、キーをBへ複製したりデータを再暗号化したりする移行は認めていない。クロスアカウントで復号を成立させるための最も適切な構成はどれか。

- A** BのGlueジョブロールにkms:Decryptを許可するIAMポリシーだけを追加する。IAMポリシーのResourceにはAのKMSキーARNを指定し、A側のキーポリシーは変更しない。
- B** AのKMSキーポリシーでBアカウントのrootプリンシパルにkms:*を許可する。B側では追加のIAMポリシーを作らず、Organizationsの同一組織所属を認可根拠にする。
- C** AのKMSキーポリシーでBのGlueジョブロールによるkms:Decryptを許可し、Bでも同じロールのIAMポリシーに対象キーARNへのkms:Decryptを許可する。必要なら暗号化コンテキスト条件も限定する。
- D** AのS3バケットポリシーにBのGlueジョブロール向けのkms:Decryptアクションを追加し、KMSキーへの認可をS3のリソースポリシーへ集約する。

ここで答えを決める

正解・解説は次のページから

ANSWER 025

問題25の正解・解説

1 正解 C

- C** AのKMSキーポリシーでBのGlueジョブロールによるkms:Decryptを許可し、Bでも同じロールのIAMポリシーに対象キーARNへのkms:Decryptを許可する。必要なら暗号化コンテキスト条件も限定する。

2 問題文の重要な条件

- KMSキーとデータはアカウントAにあり、利用ロールはアカウントBにある
- Bの特定のGlueジョブロールだけに復号を許可する
- 再暗号化やキー移行はせず、既存のカスタマーマネージドキーを使う

3 正解へ至る判断過程

1. S3の読取権限とKMSの復号権限は別評価であり、S3バケットポリシーが正しくてもKMS認可が不足すれば復号できない。
2. クロスアカウントKMS利用では、キー所有側のキーポリシーによる許可と、利用側プリンシパルのIAMポリシーによる許可の両方が必要になる。
3. 特定ロールを指定し、必要に応じて暗号化コンテキストでS3用途へ絞るCが最小権限を満たす。

4 正解が要件を満たす理由

CはKMSのクロスアカウント認可を構成する二つの側面を満たす。Aのキーポリシーが外部プリンシパルの利用を認め、BのIAMポリシーがその特定ロールに利用権限を与える。対象キーと利用文脈を限定できるため、データ再暗号化なしに最小権限を実現する。

5・6 各選択肢の検討

不適切な理由と、正解になり得る条件を対にして確認します。

A 不正解

利用側IAMの許可だけでは、キー所有アカウントAのキーポリシーがクロスアカウント利用を認めていないため不十分である。

この選択肢が正解になる条件

キーとGlueロールが同一アカウントにあり、キーポリシーがそのアカウントのIAMポリシーへ権限委任する既定の構成ならIAMポリシー追加だけで正解になり得る。

B 不正解

Bアカウント全体にkms:*を認めるのは過剰で、利用側IAMの明示許可も欠く。同じOrganizations配下であることだけではKMS利用権限にならない。

この選択肢が正解になる条件

B側で複数の管理対象ロールへ権限委任する必要がある、AのキーポリシーはBアカウントに委任し、さらにBのIAMで対象ロールだけへDecryptを限定するなら構成要素になり得る。

C 正解

所有側キーポリシーと利用側IAMの双方で特定ロールを許可し、最小権限のクロスアカウント復号を成立させる。

D 不正解

S3バケットポリシーはS3リソースへの操作を制御するもので、KMSキーに対するkms:Decryptのリソースベース認可を代替しない。

この選択肢が正解になる条件

論点がS3オブジェクトのGetObjectだけで、暗号化がSSE-S3またはKMS権限が別途整備済みなら、バケットポリシーの変更が正解になり得る。

7 今後使える判断基準

クロスアカウントのカスタマーマネージドKMSキー利用は、キー所有側のキーポリシーと利用側のIAMポリシーを一对で確認する。暗号化サービスのリソース権限とKMS権限は分けて評価する。

8 関連サービスとの比較

KMSキーポリシーはキー所有者が誰に利用を認めるかを定め、IAMポリシーは利用アカウント内のどのプリンシパルがその許可を使えるかを絞る。S3バケットポリシーはオブジェクトアクセス、KMS grantは一時的・委任的なキー利用許可に適する。

QUESTION 026

問題26

RESILIENT

本番より難しい

単一選択

Auroraのリージョン障害対策

旅行予約会社は、東京リージョンのAurora PostgreSQLに毎秒数千件の予約を記録している。Multi-AZ構成によりAZ障害には対応済みだが、リージョン全体の停止に備え、シンガポールでRPO 5分未満、RTO 10分未満を達成したい。平常時にはシンガポールの分析処理から最新に近いデータを低遅延で読みたい。DBエンジンの変更とアプリケーションの大規模改修は避け、レプリケーションサーバーの保守も行いたくない。DR演習では計画切替時のデータ損失をゼロにし、実障害時は複製遅延を確認して昇格先を選ぶ必要がある。DR側コンピュート費用は許容され、接続先変更はRunbookで自動化できる。フェイルバック手順も定期試験する。最も適切な構成はどれか。

- A** 東京のAuroraから毎時スナップショットを取得し、シンガポールへ自動コピーする。障害時に最新スナップショットから新しいクラスターを復元し、Route 53を切り替える。
- B** 東京のAuroraとシンガポールのRDS for PostgreSQL間にAWS DMSのフルロードとCDCタスクを常時実行する。障害時にはDMSタスクを停止し、RDSを手動で書き込み先にする。
- C** 東京のAuroraクラスターにリージョン内Aurora Replicaを追加し、リーダーエンドポイントをRoute 53フェイルオーバーレコードのセカンダリとして登録する。
- D** Aurora Global Databaseを作成し、シンガポールにDBインスタンスを持つセカンダリクラスターを配置する。複製遅延を監視し、計画時はスイッチオーバー、障害時はフェイルオーバー手順と接続先切替を自動化する。

ここで答えを決める

正解・解説は次のページから

ANSWER 026

問題26の正解・解説

1 正解 **D**

- D** Aurora Global Databaseを作成し、シンガポールにDBインスタンスを持つセカンダリクラスターを配置する。複製遅延を監視し、計画時はスイッチオーバー、障害時はフェイルオーバー手順と接続先切替を自動化する。

2 問題文の重要な条件

- リージョン障害に対しRPO 5分未満、RTO 10分未満が必要
- 平常時にDRリージョンで低遅延の読取りを行う
- Aurora互換性を保ち、管理レプリケーションと無損失の計画切替を求める

3 正解へ至る判断過程

1. Multi-AZやリージョン内Replicaではリージョン全体の障害に対応できず、スナップショット復元はRPOとRTOを満たしにくい。
2. Aurora Global Databaseはストレージベースのクロスリージョン複製とセカンダリのローカル読取りを提供し、既存Auroraからの変更を抑えられる。
3. 健全な計画切替では同期後のスイッチオーバー、障害時は遅延を見たフェイルオーバーを使い分けるDが運用要件まで満たす。

4 正解が要件を満たす理由

DはAuroraのまま管理されたクロスリージョン複製を利用し、セカンダリにDBインスタンスを置くことで平常時のローカル読取りと短いRTOを両立する。計画スイッチオーバーは同期してデータ損失を避けられ、非計画フェイルオーバーでは複製遅延を基に判断できる。接続先切替の自動化を含めてRTO要件に対応する。

5・6 各選択肢の検討

不適切な理由と、正解になり得る条件を対にして確認します。

A 不正解

毎時スナップショットではRPO 5分未満を満たせず、クラスター復元と起動によりRTO 10分未満も保証しにくい。平常時の読取りにも使えない。

この選択肢が正解になる条件

RPOが数時間、RTOが数時間で、DR側の平常時コンピュート費用を最小化するバックアップ&リストア戦略なら正解になり得る。

B 不正解

DMS CDCは移行や異種DB複製に有効だが、このAurora DRでは継続運用するタスク、DDL対応、切替整合性などの管理が増え、Auroraネイティブ機能より不利である。

この選択肢が正解になる条件

異種エンジンへ最小停止で移行する、または変換を伴う一方向レプリケーションが主目的ならDMSが正解になり得る。

C 不正解

リージョン内Aurora ReplicaはAZ障害と読取り拡張には有効だが、東京リージョン全体が停止すると利用できない。リーダーエンドポイントは書込み昇格済みDR先でもない。

この選択肢が正解になる条件

対象障害がインスタンスまたはAZに限定され、クロスリージョンDRが不要で読取り拡張と高速な同一リージョンフェイルオーバーが目的なら正解になり得る。

D 正解

短いRPO/RTO、DR側読取り、Aurora互換性、管理レプリケーション、計画・非計画切替の使い分けを満たす。

7 今後使える判断基準

Auroraで短いクロスリージョンRPO/RTOとDR側読取りが必要なならAurora Global Databaseを優先する。計画切替はスイッチオーバー、実障害はフェイルオーバーとして、複製遅延と接続先変更まで設計する。

8 関連サービスとの比較

RDS/Aurora Multi-AZはリージョン内HA、Aurora Replicaは同一リージョンの読取り拡張と昇格候補、Aurora Global DatabaseはクロスリージョンDRとローカル読取り向けである。スナップショットコピーは安価だがRPO/RTOが長く、DMSは主に移行・異種複製で強い。

QUESTION 027

問題27

PERFORMANCE

複合難問

単一選択

S3データレイクの走査量最適化

小売企業は5PBの購買履歴をAmazon S3にJSON形式で保存し、データ分析者がAthenaでSQLを実行している。データは毎時追加され、典型的なクエリは直近7日間、特定リージョン、全200列のうち10列未満だけを参照する。しかし現状は1回のクエリで大量のファイルと列を走査し、完了まで20分以上かかり、走査料金も増えている。既存データはS3を信頼できる原本として保持し、サーバーや常時稼働クラスターは管理したくない。元JSONは監査用に残しつつ、派生データの追加保存は許可される。クエリSQLのWHERE句は変更でき、ETLは日次で自動実行してよい。スキーマは月に数回追加され、複数分析チームで同じテーブル定義を共有したい。性能と費用を同時に改善する最も適切な構成はどれか。

- A** AWS GlueでJSONを圧縮Parquetへ変換し、クエリパターンに合う日付とリージョンでパーティション化する。Glue Data Catalogへ登録し、Athenaから必要な列とパーティションだけを問い合わせる。
- B** すべてのJSONをAmazon Redshiftのローカルテーブルへ毎時COPYし、5PB分を常時保持できるクラスターを起動する。Athenaは廃止し、全クエリをRedshiftへ移す。
- C** Amazon EMRの常時稼働クラスターでJSONをHDFSへ複製し、Hiveテーブルを各分析チームが個別に管理する。S3上の原本はバックアップ用途だけにする。
- D** JSONをそのまま維持し、S3オブジェクトを1時間ごとに1個の巨大ファイルへ結合する。Athenaのワークグループでクエリ同時実行数を増やし、全列走査を並列化する。

ここで答えを決める

正解・解説は次のページから

ANSWER 027

問題27の正解・解説

1 正解 **A**

- A** AWS GlueでJSONを圧縮Parquetへ変換し、クエリパターンに合う日付とリージョンでパーティション化する。Glue Data Catalogへ登録し、Athenaから必要な列とパーティションだけを問い合わせる。

2 問題文の重要な条件

- 典型クエリは日付・リージョンで絞り、200列中10列未満を参照する
- S3を原本とし、常時稼働クラスターを管理しない
- 共有スキーマを保ちつつ、走査時間と走査料金の両方を下げる

3 正解へ至る判断過程

1. Athenaの性能と費用は主にS3から読み取るデータ量に左右されるため、同時実行数よりファイル形式と物理配置を先に最適化する。
2. 列指向Parquetなら参照列だけを読みやすく、日付・リージョンのパーティションなら不要なオブジェクトを除外できる。圧縮も走査量を減らす。
3. Glue Data Catalogを共通メタデータにすればスキーマ共有と変更追従ができ、Athenaのサーバーレス性とS3原本を維持できる。

4 正解が要件を満たす理由

Aは列選択と絞り込み条件にデータ配置を合わせ、Parquet・圧縮・パーティションブローニングの三つで読み取り量を減らす。その結果、Athenaの実行時間と走査料金を同時に改善する。Glue Data Catalogで複数チームのメタデータを一元化し、常時クラスターを持たないという運用制約も満たす。

5・6 各選択肢の検討

不適切な理由と、正解になり得る条件を対にして確認します。

A 正解

列指向形式、適切なパーティション、共有カタログにより、設問のアクセスパターンとサーバーレス要件を同時に満たす。

B 不正解

反復的で複雑なBIにはRedshiftが有力だが、5PB全量をローカルへ常時保持すると移行・クラスター費用が大きく、S3原本・低運用という要件に対して過剰である。

この選択肢が正解になる条件

高い同時実行で複雑な結合を繰り返し、予測可能な低遅延を最優先し、ホットデータをデータウェアハウスへロードする費用を許容できるなら正解になり得る。

C 不正解

常時EMRとHDFS複製はクラスター管理とデータの二重保持を増やし、共有メタデータ要件にも反する。単純な対話型SQL最適化には運用が重い。

この選択肢が正解になる条件

長時間のSpark処理、独自ライブラリ、反復機械学習などAthenaでは実行しにくい分散計算が中心で、クラスター運用を受け入れるなら正解になり得る。

D 不正解

小ファイル削減には一部効果があるが、行指向JSONの全列読取りとパーティション不足を解消しない。同時実行数を増やしても1クエリの走査量と料金は減らない。

この選択肢が正解になる条件

既に列指向・適切なパーティション化が済み、小さすぎるファイルだけがボトルネックなら、適切なサイズへのコンパクションが正解の一部になり得る。

7 今後使える判断基準

S3上の分析では、まずクエリのフィルター列でパーティション化し、参照列が少ないならParquetやORC、圧縮、適正ファイルサイズを使う。計算資源を増やす前に走査量を減らす。

8 関連サービスとの比較

AthenaはS3を直接読むサーバーレスなアドホックSQLで、走査量が費用と性能に直結する。GlueはETLと共有カタログを担う。Redshiftは反復的な高同時実行分析、EMRは柔軟な分散処理に向くが、いずれもAthenaより構成・費用判断が増える。

QUESTION 028

問題28

SECURE

本番相当

複数選択 (2つ)

B2C認証と一時AWS資格情報

写真共有サービスは、数百万人向けのモバイルアプリを開発している。利用者はメールアドレスまたは外部ソーシャルIdPで登録・サインインし、MFAとパスワード回復を利用する。認証後、アプリはAPIを介さず利用者ごとのS3プレフィックスへ写真を直接アップロードする必要がある。端末へ長期AWSアクセスキーを埋め込むことは禁止され、利用者は自分のプレフィックス以外を読み書きできてはならない。バックエンドは受け取ったJWTでAPI認可も行う。アップロード資格情報は短期間で失効し、ユーザー停止後には新規発行を止めたい。認証基盤のサーバー運用と鍵ローテーションもAWSへ委ねる。運用チームは利用者ごとのIAMユーザー作成を避けた。要件を満たすために採用すべきコンポーネントを2つ選択してください。

- A** Amazon Cognitoユーザープールを利用者ディレクトリとIdP連携に使い、サインアップ、サインイン、MFAを処理してJWTをアプリへ発行する。
- B** Cognito IDプールをユーザープールと連携し、認証済み利用者へIAMロールに基づく一時AWS資格情報を発行する。ポリシー変数などでS3プレフィックスを利用者単位に制限する。
- C** 利用者ごとにIAMユーザーとアクセスキーを作成し、モバイルOSの安全なストレージへ保存する。外部IdPのユーザー名をIAMユーザー名に対応させる。
- D** ユーザープールのIDトークンをS3 PutObject APIのAuthorizationヘッダーへ直接設定する。S3バケットポリシーではJWTのsubクレームを評価してプレフィックスを制限する。

ここで答えを決める

正解・解説は次のページから

ANSWER 028

問題28の正解・解説

1 正解 **A・B**

- A** Amazon Cognitoユーザープールを利用者ディレクトリとIdP連携に使い、サインアップ、サインイン、MFAを処理してJWTをアプリへ発行する。
- B** Cognito IDプールをユーザープールと連携し、認証済み利用者へIAMロールに基づく一時AWS資格情報を発行する。ポリシー変数などでS3プレフィックスを利用者単位に制限する。

2 問題文の重要な条件

- B2C利用者の登録、外部IdP、MFA、JWT発行が必要
- モバイルアプリからS3へ直接アクセスするが長期AWSキーは禁止
- 利用者ごとにS3プレフィックスを最小権限で分離する

3 正解へ至る判断過程

1. アプリ利用者の認証と、AWS APIを署名する資格情報の発行は別の役割として分ける。
2. ユーザープールは利用者認証とJWTを、IDプールは認証トークンと引き換えにIAMロール由来の一時AWS資格情報を提供する。
3. IAMロールのポリシーを利用者識別子に対応するプレフィックスへ限定すれば、長期キーなしの直接S3アクセスを実現できる。

4 正解が要件を満たす理由

AでB2Cのサインアップ、ソーシャル連携、MFA、JWTを扱い、Bでその認証結果をAWSの短期資格情報へ交換する。ユーザーごとのIAMユーザーは不要で、IDプールの認証済みロールとS3ポリシーによりプレフィックス分離できる。APIはユーザープールJWT、S3は署名済みAWS APIという二つの要件を適切に分担する。

5・6 各選択肢の検討

不適切な理由と、正解になり得る条件を対にして確認します。

A 正解

利用者ディレクトリ、フェデレーション、MFA、JWTというアプリ認証要件を担当する。

B 正解

ユーザープールの認証結果から期限付きAWS資格情報を発行し、IAMでS3アクセスを利用者単位に絞れる。

C 不正解

数百万人分のIAMユーザー管理と端末上の長期アクセスキーは、明示されたセキュリティ・運用要件に反する。IAMユーザーは一般消費者ID管理向けではない。

この選択肢が正解になる条件

少人数の人間管理者がAWSコンソールやCLIを使用し、フェデレーションを採用できず、厳格なキー管理とローテーションを行う限定環境ならIAMユーザーが必要になる場合がある。

D 不正解

CognitoのJWTはそのままS3 APIのSigV4用AWS資格情報にはならず、S3バケットポリシーが任意のJWTクレームを直接検証する構成でもない。

この選択肢が正解になる条件

S3へ直接アクセスせず、API Gatewayや独自APIがJWTを検証してサーバー側でS3操作を代行する構成なら、JWTをAPI認可に用いる方法が正解になり得る。

7 今後使える判断基準

Cognitoではユーザープールを『誰かを認証してJWTを出すもの』、IDプールを『認証結果から一時AWS資格情報を出すもの』として区別する。AWSサービスへのクライアント直接アクセスがあるときだけ後者を追加する。

8 関連サービスとの比較

ユーザープールはB2C認証、ユーザーディレクトリ、MFA、OAuth/OIDCトークンを担当する。IDプールはユーザープールや外部IdPのトークンをIAMロール由来の一時資格情報へ交換する。IAM Identity Centerは従業員のAWSアクセス、IAMユーザーは長期AWS主体向けで、一般消費者認証とは用途が異なる。

QUESTION 029

問題29

RESILIENT

本番より難しい

単一選択

S3複製時間と削除保護

報道会社は東京のS3バケットへ記事素材を保存し、リージョン災害時には大阪のバケットから配信する。新規オブジェクトの99.99%を15分以内に複製するSLAが必要で、複製遅延を監視したい。東京側でアプリケーションが誤って通常のDELETEを実行しても、大阪側の直近バージョンは直ちに失われてはならない。既存の20TBも初回移行対象である。両バケットは同じ会社の別アカウントにあり、今後の新規オブジェクトと既存オブジェクトを管理サービスで複製したい。オブジェクト名は更新時に再利用されるため、過去版からの復旧も必要である。独自コピーLambdaの再試行や容量管理は避け、複製失敗を通知できるようにする。要件を最もよく満たす構成はどれか。

- A** 両バケットのバージョニングを無効にしたままS3 Transfer Accelerationを有効にし、東京のPutObjectイベントごとにLambdaで大阪へ同期コピーする。DELETEイベントも同じ関数で即時反映する。
- B** 両バケットでバージョニングを有効にし、CRRにS3 Replication Time Controlと複製メトリクスを設定する。削除マーカー複製は無効にし、既存オブジェクトにはS3 Batch Replicationを実行する。
- C** S3 Multi-Region Access Pointだけを作成し、東京と大阪のバケットを登録する。レプリケーションルールは作成せず、アクセスポイントが読み書き時に自動でオブジェクトを同期すると想定する。
- D** 両バケットでバージョニングと標準CRRを有効にし、削除マーカー複製も有効にする。15分要件はAmazon CloudWatchの定期ポーリングと運用担当者の手動再送で保証する。

ここで答えを決める

正解・解説は次のページから

ANSWER 029

問題29の正解・解説

1 正解 **B**

- B** 両バケットでバージョニングを有効にし、CRRにS3 Replication Time Controlと複製メトリクスを設定する。削除マーカー複製は無効にし、既存オブジェクトにはS3 Batch Replicationを実行する。

2 問題文の重要な条件

- 新規オブジェクトの99.99%を15分以内に複製するSLAと監視が必要
- ソースの通常DELETEを直ちに宛先へ伝播させず、復旧可能にする
- 既存20TBと今後の新規データの両方を管理サービスで複製する

3 正解へ至る判断過程

1. クロスリージョンライブ複製には両バケットのバージョニングが前提で、予測可能な15分の複製にはS3 RTCを選ぶ。
2. 通常DELETEはバージョニング管理バケットに削除マーカーを作るため、その複製を無効にすれば宛先の現行表示を誤削除から分離できる。
3. ライブレプリケーションは設定後の新規オブジェクトが中心なので、既存分はS3 Batch Replicationで補う。

4 正解が要件を満たす理由

Bはバージョニング、CRR、S3 RTC、メトリクスにより、新規オブジェクトのクロスリージョン複製と15分SLA・可視化を満たす。削除マーカーを複製しないためソースの通常DELETEが宛先の現行オブジェクトを隠さず、既存20TBはBatch Replicationで同じ管理基盤へ取り込める。

5・6 各選択肢の検討

不適切な理由と、正解になり得る条件を対にして確認します。

A 不正解

Transfer Accelerationはクライアント転送を高速化する機能でレプリケーションSLAではない。独自Lambda同期は容量、再試行、順序、監視の運用を増やし、DELETE伝播も保護要件に反する。

この選択肢が正解になる条件

少量データに独自変換を加えて複製し、複製時間SLAがなく、削除も両側で即時一致させる必要があるならイベント駆動コピーが候補になり得る。

B 正解

RTCのSLA、遅延監視、削除の分離、既存データの一括複製をすべて満たす。

C 不正解

Multi-Region Access Pointは複数リージョンのバケットへのグローバルアクセス経路を提供するが、登録だけでデータを自動複製しない。別途レプリケーション設定が必要である。

この選択肢が正解になる条件

双方向レプリケーションが既に構成済みで、利用者を最も近い正常なバケットヘルレーティングする単一グローバルエンドポイントが必要なら正解の一部になり得る。

D 不正解

標準CRRには設問の15分SLAがなく、手動再送でサービスレベルを保証できない。さらに削除マーカを複製するとソースの誤DELETEが宛先にも反映される。

この選択肢が正解になる条件

予測可能な複製時間が不要で、ソースと宛先の論理削除状態を一致させることが要件なら標準CRRと削除マーカ複製が正解になり得る。

7 今後使える判断基準

S3複製問題では、新規か既存か、複製時間SLAの有無、削除マーカを伝播するかを別々に判断する。RTCは時間、Batch Replicationは既存分、削除マーカ設定は誤削除隔離を決める。

8 関連サービスとの比較

標準CRRは管理された非同期クロスリージョン複製、S3 RTCはその複製時間SLAと可視化を追加する。S3 Batch Replicationは既存オブジェクトの一括複製、Multi-Region Access Pointはアクセスルーティングであり、データ複製そのものではない。

QUESTION 030

問題30

PERFORMANCE

複合難問

単一選択

RedshiftとS3履歴データの統合分析

通信会社は、直近12か月の請求・顧客データ12TBをAmazon Redshiftで分析し、BIダッシュボードから複雑な結合を高頻度に行っている。5年以上前までのネットワーク履歴300TBはS3上のParquetに日付と地域でパーティション化され、月数回の調査時だけ参照される。新要件では、Redshiftの顧客テーブルとS3の履歴を1つのSQLで結合したい。300TBを常時ウェアハウスへロードする費用は避け、既存BIツールとRedshiftの権限管理を維持したい。履歴のみの単発調査より、現行データとの結合性能と運用の一貫性を優先する。S3とRedshiftは同一リージョンにあり、Glueカタログ作成は許可される。分析者は日付条件を指定でき、元データの移動は避けたい。最も適切な構成はどれか。

- A** Redshiftを廃止し、直近12TBもParquetとしてS3へアップロードする。すべてをAthenaで問い合わせ、BIの接続先と権限モデルを全面的に変更する。
- B** 履歴300TBをすべてRedshiftのローカルテーブルへCOPYし、最大需要に合わせてクラスターを常時拡張する。S3の原本は削除して二重管理を避ける。
- C** Glue Data Catalogに履歴テーブルを登録し、Redshift Spectrumの外部スキーマからS3を参照する。現行データはRedshiftローカルに維持し、パーティション条件を含むSQLで両者を結合する。
- D** 月次でAmazon EMRクラスターを起動し、S3履歴とRedshiftスナップショットをHDFSへコピーしてSparkで結合する。結果だけを新しいRedshiftクラスターへロードする。

ここで答えを決める

正解・解説は次のページから

ANSWER 030

問題30の正解・解説

1 正解 C

- C** Glue Data Catalogに履歴テーブルを登録し、Redshift Spectrumの外部スキーマからS3を参照する。現行データはRedshiftローカルに維持し、パーティション条件を含むSQLで両者を結合する。

2 問題文の重要な条件

- 高頻度・複雑な現行分析は既存Redshiftで良好に動作している
- 300TBの履歴は低頻度参照で、全量ロード費用を避ける
- RedshiftのSQL、BI接続、権限を維持してローカル表とS3を結合する

3 正解へ至る判断過程

1. 高頻度のホットデータと低頻度のコールド履歴では、同じ保存方式に統一するよりアクセス頻度に応じて配置を分ける。
2. Redshift SpectrumはRedshiftからS3外部テーブルを直接問い合わせ、ローカルテーブルと同じSQLで結合できる。
3. 既存RedshiftとBIを維持しつつ、Parquetとパーティションで履歴走査を抑えるCが性能・費用・移行制約の均衡を取る。

4 正解が要件を満たす理由

Cは頻繁に使う12TBをRedshiftローカルに残して既存の結合性能とBI運用を保ち、低頻度の300TBだけをSpectrumでS3から必要時に読む。外部スキーマと共有カタログにより1つのSQLで結合でき、履歴全量のロード・常時コンピュート・データ重複を避けられる。

5・6 各選択肢の検討

不適切な理由と、正解になり得る条件を対にして確認します。

A 不正解

Athenaは単発のS3分析には適するが、既にRedshiftで最適化された高頻度の複雑結合、BI接続、権限を全面移行するため、設問の優先事項に反する。

この選択肢が正解になる条件

すべての分析が低頻度のアドホックSQLで、常時ウェアハウスが不要、BIと権限の移行を許容できるならAthena中心が正解になり得る。

B 不正解

履歴全量をロードすれば性能は得やすいが、月数回しか使わない300TBの保存・コンピュータ費用が過大で、全量ロード回避という明示条件を満たさない。原本削除もデータレイク活用を損なう。

この選択肢が正解になる条件

履歴全体を毎日反復して複雑に結合し、予測可能な低遅延が費用より重要で、Redshift内へ保持する合理性があるなら正解になり得る。

C 正解

ホットデータのRedshift性能と、コールドデータのS3経済性を、Redshift SQL内で統合する。

D 不正解

EMRは複雑なSpark処理には適するが、データコピー、クラスター起動、別Redshiftへの結果ロードが増え、対話的な既存BIからの一体的SQLという要件を満たしにくい。

この選択肢が正解になる条件

SQLでは表現しにくい大規模な機械学習前処理やカスタムSpark処理を月次バッチで行い、結果だけを提供すればよいなら正解になり得る。

7 今後使える判断基準

Redshiftのホット表とS3の低頻度大容量データを同じSQLで結合するならRedshift Spectrumを検討する。アドホックなS3単独分析はAthena、反復BIはRedshiftローカル、任意分散処理はEMRと切り分ける。

8 関連サービスとの比較

Redshiftローカルテーブルは反復的な低遅延分析、SpectrumはRedshiftからS3外部表を結合する用途、AthenaはクラスターなしのS3アドホックSQL、EMRはSparkなどの柔軟な大規模処理に向く。SpectrumとAthenaはいずれも列指向形式とパーティションが重要である。

QUESTION 031

問題31

SECURE

本番相当

単一選択

集中型egressのステートフル検査

企業はAWS Organizations内の20アカウントで40個のスポークVPCを運用している。各VPCのプライベートサブネットからのインターネット向け通信を、ネットワークアカウントで一元的に許可リスト検査し、ログを保管したい。スポークにInternet GatewayやNAT Gatewayは配置せず、検査後の許可トラフィックだけを共有egress VPCのNAT Gatewayから送信する。TCPセッションをステートフルに検査するため、往路と復路が同じファイアウォール経路を通る必要がある。各スポークCIDRは重複せず、中央ハブ接続を許可している。検査ポリシーは全アカウントで共通化し、独自ファイアウォールOSの保守は避けたい。AZ障害時にも別AZで通信を継続し、新しいVPC追加時のピアリング作業を最小化したい。最も適切な構成はどれか。

- A** すべてのスポークVPCをフルメッシュのVPCピアリングでegress VPCへ接続する。egress VPCの1つのパブリックサブネットにNATインスタンスとホスト型IDSを配置する。
- B** 各スポークVPCの各AZにNAT GatewayとAWS Network Firewallを個別配置し、Firewall Managerを使わず各アカウントがルールを管理する。
- C** AWS PrivateLinkのインターフェイスエンドポイントサービスとしてNAT Gatewayを公開し、各スポークのデフォルトルートをそのインターフェイスエンドポイントへ設定する。
- D** Transit Gatewayでスポークと集中検査VPCを接続し、検査VPCの複数AZにNetwork Firewallエンドポイントを置く。アプライアンスモードとTGW/VPCルートで双方向を同じ検査経路へ通し、その後egress VPCのNAT Gatewayへ送る。

ここで答えを決める

正解・解説は次のページから

ANSWER 031

問題31の正解・解説

1 正解 **D**

- D** Transit Gatewayでスポークと集中検査VPCを接続し、検査VPCの複数AZにNetwork Firewallエンドポイントを置く。アプライアンスモードとTGW/VPCルートで双方向を同じ検査経路へ通し、その後egress VPCのNAT Gatewayへ送る。

2 問題文の重要な条件

- 多数のアカウント・VPCをハブ型で接続し、追加作業を抑える
- 全インターネットegressを一元検査してから共有NATへ送る
- ステートフル検査の経路対称性とMulti-AZ継続性が必要

3 正解へ至る判断過程

1. 多数VPCの推移的なハブ接続にはTransit Gatewayが適し、フルメッシュピアリングの接続数とルート管理を避けられる。
2. Network Firewallを複数AZへ配置して集中検査し、許可トラフィックだけをNATへ送ること
でセキュリティとegressを分離できる。
3. ステートフル装置では非対称ルーティングを避ける必要があるため、検査VPCアタッチメン
トのアプライアンスモードと往復ルートを設計するDを選ぶ。

4 正解が要件を満たす理由

DはTGWによる大規模ハブ接続、Network Firewallによる一元的なL3～L7検査、集中NATによる外向き通信を組み合わせる。複数AZのエンドポイントとアプライアンスモードを使い、往路と復路を同じAZのステートフル検査経路へ固定することで、可用性とセッション整合性も満たす。

5・6 各選択肢の検討

不適切な理由と、正解になり得る条件を対にして確認します。

A 不正解

VPCピアリングは推移的ルーティングを提供せず、多数VPCのフルメッシュ管理が増える。単一AZのNATインスタンスとIDSも可用性、スケーリング、管理負荷で要件に劣る。

この選択肢が正解になる条件

接続VPCが2〜3個に限定され、推移ルーティング不要で、特定の仮想アプライアンス機能を自ら運用する必要があるならピアリング構成が候補になり得る。

B 不正解

各VPC配置は通信の局所性とAZ耐性には有効だが、40 VPC分のNAT・Firewall費用とルール運用が増え、『一元化』とコスト・運用要件を満たさない。

この選択肢が正解になる条件

各VPCが独立したセキュリティ境界で、集中経路の障害半径やTGW転送料を避けることを最優先し、重複費用を許容するなら正解になり得る。

C 不正解

PrivateLinkは特定サービスを非公開公開する仕組みで、NAT Gatewayをエンドポイントサービス化して任意のインターネットegressや推移ルーティングを提供する構成ではない。

この選択肢が正解になる条件

複数VPCから特定の社内TCPサービスだけを一方向に公開し、アドレス重複や推移接続を避けたいならPrivateLinkが正解になり得る。

D 正解

多数VPCの集約、集中egress検査、対称なステートフル経路、Multi-AZを一つの管理可能な構成で満たす。

7 今後使える判断基準

多数VPCの集中検査では、接続集約にTGW、検査にNetwork Firewall、外向き変換にNATを割り当てる。ステートフル装置をTGW経由に置くときはアプライアンスモードと往復ルート対称性を必ず確認する。

8 関連サービスとの比較

Transit Gatewayは多数VPCの推移的ハブ接続、VPCピアリングは少数VPCの直接接続、PrivateLinkは特定サービスの限定公開に向く。Network Firewallは管理型ステートフル検査、NAT Gatewayは外向きアドレス変換であり、役割は置き換え関係ではない。

QUESTION 032

問題32

RESILIENT

本番より難しい

複数選択 (2つ)

EFS RegionalとOne Zoneの使い分け

企業はNFS互換の共有ストレージを必要とする2つの新規システムを設計している。システムXは3つのAZにまたがるECSサービスで、利用者が作成した文書を保持する。単一AZ喪失時もデータアクセスを継続し、アプリケーションを別AZで動かす必要がある。システムYは1つのAZに固定された一時的な機械学習前処理で、入力S3にあり、出力も完了時にS3へ保存される。Yの作業領域は消失しても6時間で再生成でき、保存費用を最優先する。XのRPOはほぼゼロ、YのRPOは作業単位全損を許容する。両者ともPOSIXファイル操作と複数タスクからの同時マウントが必要で、運用チームはファイルサーバーを管理したくない。最適な設計を2つ選択してください。

- A** システムXにはEFS Regionalファイルシステムを使用し、ECSが利用する各AZにマウントターゲットを作成する。タスクは同一AZのマウントターゲットへ接続する。
- B** システムXにはEFS One Zoneを1つ作り、3つのAZにマウントターゲットを追加する。AZ障害時はDNSが別AZの同じデータへ自動で切り替える。
- C** システムYにはEFS One Zoneを処理タスクと同じAZに作成する。作業領域は再生成可能とし、必要ならS3原本と完了済み出力だけを保護する。
- D** システムYにはEFS Regionalを使用し、3AZすべてにマウントターゲットを作成する。タスクは1AZ固定のままだが、作業領域にもAZ障害耐性を持たせる。

ここで答えを決める

正解・解説は次のページから

ANSWER 032

問題32の正解・解説

1 正解 **A・C**

- A** システムXにはEFS Regionalファイルシステムを使用し、ECSが利用する各AZにマウントターゲットを作成する。タスクは同一AZのマウントターゲットへ接続する。
- C** システムYにはEFS One Zoneを処理タスクと同じAZに作成する。作業領域は再生成可能とし、必要ならS3原本と完了済み出力だけを保護する。

2 問題文の重要な条件

- Xは利用者データを保持し、単一AZ障害時も別AZから継続アクセスする
- Yは単一AZ固定で作業データをS3から再生成でき、費用最優先である
- 両システムとも管理不要のNFS/POSIX共有ストレージが必要

3 正解へ至る判断過程

1. データ自体にMulti-AZ耐性が必要か、計算と同じ単一AZで十分かをシステムごとに分ける。
2. EFS Regionalは複数AZにデータを冗長化し、各AZのマウントターゲットでXの継続性を支える。
3. EFS One Zoneはデータを単一AZ内に保存する代わりに低コストであり、再生成可能かつ単一AZ固定のYに合う。

4 正解が要件を満たす理由

AはXのデータを複数AZに冗長化し、各AZのタスクがローカルなマウントターゲットを使えるためAZ障害時の継続性を満たす。CはYの計算とファイルシステムを同じAZに置き、再生成可能なスクラッチ領域へOne Zoneのコスト優位を適用する。可用性価値の異なるデータを同じ高価な構成へ揃えない点が最適である。

5・6 各選択肢の検討

不適切な理由と、正解になり得る条件を対にして確認します。

A 正解

複数AZに冗長化されたデータとAZごとのアクセス経路により、XのAZ障害耐性を実現する。

B 不正解

EFS One Zoneは単一AZにデータを保存し、作成できるマウントターゲットもそのAZの1つである。3AZへ同一データを冗長化する説明は成立しない。

この選択肢が正解になる条件

Xが同じ1AZに固定され、データをバックアップや原本から再生成でき、AZ障害時の継続が不要ならOne Zoneが正解になり得る。

C 正解

Yの単一AZ配置、再生成可能性、費用最優先という条件にOne Zoneを合わせている。

D 不正解

RegionalはYにも機能上利用できるが、タスク自体が単一AZ固定で作業領域も再生成可能なため、追加のMulti-AZ耐性に費用を払う価値が要件上ない。

この選択肢が正解になる条件

Yの作業途中状態を失うとRPO/RTOを満たせず、計算タスクも障害時に別AZへ移動できるならRegionalが正解になり得る。

7 今後使える判断基準

EFSタイプは『共有できるか』ではなく『データがAZ喪失に耐える必要があるか』で決める。再生成不能・Multi-AZ計算はRegional、再生成可能・単一AZ固定でコスト優先ならOne Zoneを検討する。

8 関連サービスとの比較

EFS Regionalはデータを複数AZに保存しAZごとにマウントターゲットを作れる。EFS One Zoneは単一AZ保存・単一マウントターゲットで低コストだがAZ喪失リスクを受け入れる。どちらもNFS/POSIXを提供し、S3はオブジェクト原本や再生成元に向く。

QUESTION 033

問題33

SECURE

複合難問

単一選択

変更不能な長期保管とランサムウェア対策

証券会社は、約定記録を作成後7年間、WORMとしてAmazon S3に保持しなければならない。規制では、保持期間中はAWSアカウントのrootユーザーを含むいかなる管理者も、オブジェクトを削除、上書き、保持期間短縮できないことが求められる。ランサムウェアで本番アカウントの管理権限が奪われても保護を維持するため、別のログ保管アカウントにもコピーする。アプリケーションは同じキー名で訂正版を追加できるが、以前の版も残す必要がある。7年経過後の削除はLifecycleで自動化してよく、保持中の読取りは通常のS3 APIで行いたい。暗号化だけでは改ざん防止要件を満たさないと監査で指摘された。法務部門による例外削除や管理者のバイパスは認められない。要件を最もよく満たす構成はどれか。

- A 両バケットでS3バージョニングとObject Lockを有効化し、既定保持をCompliance modeで7年にする。必要権限を持つレプリケーションロールで別アカウントへ複製し、宛先でも保持設定を維持する。
- B 両バケットでS3バージョニングとObject Lock Governance modeを有効にし、7年保持を設定する。セキュリティ管理者にはs3:BypassGovernanceRetentionを許可する。
- C S3バージョニングとMFA Deleteを有効にし、rootユーザーのMFAデバイスを金庫に保管する。ライフサイクルで7年後に旧バージョンを削除する。
- D 毎日のS3 Inventoryを生成し、オブジェクトをS3 Glacier Deep Archiveへ移行する。バケットポリシーでDeleteObjectを拒否し、緊急時にはrootユーザーがポリシーを変更できるようにする。

ここで答えを決める

正解・解説は次のページから

ANSWER 033

問題33の正解・解説

1 正解 **A**

- A** 両バケットでS3バージョニングとObject Lockを有効化し、既定保持をCompliance modeで7年にする。必要権限を持つレプリケーションロールで別アカウントへ複製し、宛先でも保持設定を維持する。

2 問題文の重要な条件

- 7年間のWORMで、rootを含め誰も削除・上書き・短縮できない
- 同じキー名の訂正版を新バージョンとして追加し、旧版も保持する
- 本番アカウント侵害に備え、別アカウントのコピーにも不変性を持たせる

3 正解へ至る判断過程

1. 単なるアクセス拒否ではなく、特権管理者でも解除できない保持制御が必要なためObject Lock Compliance modeを選ぶ。
2. Object Lockはオブジェクトバージョン単位で働くため、バージョニングにより訂正版を追加しながら過去版を保持できる。
3. 別アカウントへ複製し、宛先側にも保持を適用することで、認証情報侵害とアカウント境界をまたぐ防御を強化する。

4 正解が要件を満たす理由

AのCompliance modeでは保持中のオブジェクトバージョンをrootユーザーでも削除できず、保持モード変更や期間短縮もできない。バージョニングにより同じキーの訂正版は新しい版として保存され、旧版はロックされたまま残る。別アカウントへの複製と宛先保持を組み合わせ、ランサムウェアの障害半径も縮小する。

5・6 各選択肢の検討

不適切な理由と、正解になり得る条件を対にして確認します。

A 正解

Compliance mode、バージョニング、別アカウント複製により、管理者バイパスなしの7年WORMと版管理を満たす。

B 不正解

Governance modeは特別な権限を持つ主体が保持を回避できる。BypassGovernanceRetentionを管理者へ与える構成は、例外もrootバイパスも許さない条件に反する。

この選択肢が正解になる条件

運用管理者による承認済みの保持短縮・例外削除が必要で、強力な権限統制と監査でバイパスを許容するならGovernance modeが正解になり得る。

C 不正解

バージョニングとMFA Deleteは誤削除耐性を高めるが、規制上のWORM保持を提供せず、rootを含む全主体が7年間変更不能という要件をObject Lockのように強制しない。

この選択肢が正解になる条件

規制WORMではなく、管理者による意図的な復旧操作を残しつつ偶発削除への防御を強めるだけなら正解の一部になり得る。

D 不正解

Deep Archiveは低コストの保管クラスであり、それ自体は削除不能制御ではない。変更可能なバケットポリシーはroot侵害時に解除され、WORM要件を満たさない。

この選択肢が正解になる条件

目的が不変性ではなく、年単位でほぼアクセスしないデータの保管費削減で、削除防止を別のObject Lock設定で担保するならDeep Archive移行が正解になり得る。

7 今後使える判断基準

特権者を含めた変更不能性が要件ならObject Lock Compliance mode、承認済み例外を残すならGovernance modeを選ぶ。保存クラス、バケットポリシー、バージョニングだけをWORMの代替にしない。

8 関連サービスとの比較

Object Lock Complianceは保持中に誰も解除できないWORM、Governanceはバイパス権限を持つ管理者が例外操作できる。S3 Versioningは版の復元、MFA Deleteは削除操作の追加認証、Glacier Deep Archiveは低コスト保管であり、それぞれ不変保持とは役割が異なる。

QUESTION 034

問題34

COST

本番相当

単一選択

NAT Gateway経由のAWSサービス通信削減

分析会社は3つのAZのプライベートサブネットに200台のEC2を稼働し、同一リージョンのAmazon S3とDynamoDBへ毎月300TBアクセスしている。現在は全通信が1つのAZにあるNAT Gatewayを経由し、NATの処理料金とAZ間データ転送料が高額になっている。EC2は外部パッケージ取得のため限定的なインターネット送信も引き続き必要である。S3とDynamoDBへの通信はパブリックインターネットを経由させず、バケットおよびテーブルへの最小権限も維持したい。通信の70%以上がこの2サービス向けで、オンプレミスやピアVPCから同じエンドポイントへ到達する要件はない。ルートテーブルの変更は許可される。アプリケーションコードやDNS名は変更せず、最も低コストに改善する構成はどれか。

- A** 各AZにS3とDynamoDBのインターフェイスVPCエンドポイントを作成し、Private DNSを有効にする。既存NAT Gatewayはすべての通信に引き続き使用する。
- B** S3とDynamoDBのゲートウェイVPCエンドポイントを作成し、3AZのプライベートサブネットのルートテーブルへ関連付ける。エンドポイントポリシーを限定し、NATはその他のインターネット通信だけに残す。
- C** 単一NAT Gatewayを廃止して各AZにNAT Gatewayを配置し、S3とDynamoDBを含む全トラフィックを同じAZのNATへ送る。NAT障害時は別AZへ自動迂回させる。
- D** EC2へパブリックIPv4アドレスを割り当て、Internet Gateway経由でS3とDynamoDBへ直接接続する。セキュリティグループの送信先を各サービスのIP範囲に限定する。

ここで答えを決める

正解・解説は次のページから

ANSWER 034

問題34の正解・解説

1 正解 **B**

- B** S3とDynamoDBのゲートウェイVPCエンドポイントを作成し、3AZのプライベートサブネットのルートテーブルへ関連付ける。エンドポイントポリシーを限定し、NATはその他のインターネット通信だけに残す。

2 問題文の重要な条件

- 同一リージョンのS3・DynamoDB通信300TBがNAT費用の主因
- アプリケーションやDNSを変えず、非インターネット経路と最小権限を求める
- 限定的な一般インターネット送信は残す必要がある

3 正解へ至る判断過程

1. S3とDynamoDBには追加料金のないゲートウェイVPCエンドポイントがあり、対象サービスのプレフィックスへのルートはNATより具体的にできる。
2. 各プライベートサブネットのルートテーブルへ関連付ければ、AZをまたいで単一NATへ送っていた大量通信とその処理料金を回避できる。
3. その他の0.0.0.0/0通信はNATへ残し、エンドポイントポリシーとS3/DynamoDB側ポリシーでアクセスを限定する。

4 正解が要件を満たす理由

Bは大量のS3・DynamoDB通信だけを無料のゲートウェイエンドポイントへ分岐し、NAT GatewayのGB処理料金と単一AZ NATまでのAZ間転送を削減する。既存サービスDNSとSDKを維持でき、通信はAWSネットワーク内に留まる。NATは外部パッケージ取得用に残るため機能要件も失わず、エンドポイントポリシーで最小権限を補強できる。

5・6 各選択肢の検討

不適切な理由と、正解になり得る条件を対にして確認します。

A 不正解

インターフェイスエンドポイントでも私設接続は可能だが、AZごとの時間料金とデータ処理料金が発生する。同一VPCからS3/DynamoDBへ大量アクセスする設問では無料のゲートウェイ型より高くなりやすい。

この選択肢が正解になる条件

オンプレミス、別リージョンのピアVPC、またはTransit Gateway経由からプライベートIPでS3/DynamoDBへ到達する必要がある、ゲートウェイ型の到達範囲では足りないなら正解になり得る。

B 正解

S3/DynamoDB向け大量通信を追加料金なしでNATから外し、既存コードと一般インターネット経路を保てる。

C 不正解

AZごとのNATは可用性とAZ間転送の改善にはなるが、300TBのNAT処理料金自体は残り、NATの時間料金も増える。AWSサービスへの私設経路という条件にも最適ではない。

この選択肢が正解になる条件

主な通信先がVPCエンドポイント非対応のインターネットで、AZ障害耐性とクロスAZ NAT経路の解消が目的なら各AZのNAT Gatewayが正解になり得る。

D 不正解

プライベートEC2を公開アドレス化すると攻撃面とIPv4費用が増え、パブリックインターネットを経由させない条件に反する。IP範囲管理も運用負荷が高い。

この選択肢が正解になる条件

インスタンス自体が外部から直接到達可能であることが必要な公開ワークロードで、適切な入口制御を行うならパブリックサブネットが選択肢になり得る。

7 今後使える判断基準

NAT費用問題では通信先別のバイト量を確認し、S3/DynamoDBならまずゲートウェイエンドポイントでNAT経路から外す。一般インターネット用NATのAZ配置最適化は、その後に残る通信について判断する。

8 関連サービスとの比較

S3/DynamoDBゲートウェイエンドポイントは追加料金なしでルートテーブルに組み込む。インターフェイスエンドポイントはPrivateLinkのENIとして時間・処理料金がかかるが、オンプレミスなどからの到達性に利点がある。NAT Gatewayは一般外向き通信に必要なだが時間・GB処理料金が発生する。

QUESTION 035

問題35

SECURE

本番より難しい

単一選択

OACで保護するSSE-KMS暗号化S3オリジン

動画配信会社は、ap-northeast-1のS3バケットに会員向け教材を保存し、CloudFrontから配信している。すべてのオブジェクトは監査部門が管理するカスタマーマネージドKMSキーでSSE-KMS暗号化される。現在は一部オブジェクトをS3の公開URLでも取得でき、退会後の利用者がCloudFrontを経由せずアクセスできることが判明した。S3 Block Public Accessは有効化し、取得経路を特定のCloudFrontディストリビューションだけに限定したい。オリジンへの通信にもHTTPSを使用し、アプリケーションで署名URLを毎回発行する運用は増やしたくない。既存の暗号化方式とキャッシュ配信は維持する必要がある。これらの要件を最も安全かつ運用負荷を低く満たす構成はどれか。

- A** S3静的ウェブサイトエンドポイントをCloudFrontオリジンにし、バケットポリシーで全利用者のGetObjectを許可する。CloudFrontのビューワプロトコルポリシーだけをHTTPSにする。
- B** S3 RESTエンドポイントにOrigin Access Identity (OAI) を設定する。KMSキーポリシーは変更せず、S3バケットポリシーでOAIのCanonical UserにGetObjectを許可する。
- C** S3 RESTエンドポイントにOrigin Access Control (OAC) を設定してリクエストへ常に署名する。CloudFrontサービスプリンシパルをディストリビューションARNで限定してバケットポリシーとKMSキーポリシーの両方に許可する。
- D** アプリケーションが各オブジェクトのS3 presigned URLを発行し、そのURLへ利用者をリダイレクトする。バケットは非公開にし、CloudFrontは動的ページだけに使用する。

ここで答えを決める

正解・解説は次のページから

ANSWER 035

問題35の正解・解説

1 正解 C

- C** S3 RESTエンドポイントにOrigin Access Control (OAC) を設定してリクエストへ常に署名する。CloudFrontサービスプリンシパルをディストリビューションARNで限定してバケットポリシーとKMSキーポリシーの両方に許可する。

2 問題文の重要な条件

- S3オブジェクトを特定のCloudFrontディストリビューション経由だけで取得させる
- オブジェクトはカスタマーマネージドKMSキーによるSSE-KMSのまま維持する
- S3 Block Public Access、HTTPS、CloudFrontキャッシュを維持し、署名URL発行の運用を増やさない

3 正解へ至る判断過程

1. S3オリジンを非公開化するには、公開ウェブサイトエンドポイントではなくS3 RESTエンドポイントとCloudFrontのオリジン認証を組み合わせる。
2. SSE-KMSオブジェクトの復号にはS3のGetObject許可だけでなく、CloudFrontからのオリジン要求を処理できるKMSキーポリシーも必要になる。
3. OACはSigV4署名、SSE-KMS、ディストリビューションARNによる条件付けに対応するため、経路制限と低運用負荷を同時に満たす。

4 正解が要件を満たす理由

Cは、OACでCloudFrontからS3への要求を署名し、バケットポリシーのPrincipalをCloudFrontサービス、aws:SourceArnを対象ディストリビューションに限定できる。さらに同じ限定条件を考慮してKMSキーの復号権限を付与するため、S3とKMSの二つの認可面を満たす。S3は公開不要で、SSE-KMSとCloudFrontキャッシュも維持でき、アプリケーションによるURL発行処理も不要である。

5・6 各選択肢の検討

不適切な理由と、正解になり得る条件を対にして確認します。

A 不正解

S3静的ウェブサイトエンドポイントはOACで保護できず、全利用者にGetObjectを許可すると直接アクセスを防げない。CloudFront側のHTTPS設定もCloudFrontから公開S3への迂回を封じない。

この選択肢が正解になる条件

教材が機密ではなくS3からの直接公開を許容し、SSE-KMSではなく公開ウェブサイトとして配信することが要件なら選択肢になり得る。

B 不正解

OAIは従来の非公開S3オリジンには使えるが、この構成ではSSE-KMSキーを利用するための許可が欠ける。OACが推奨され、署名方式やポリシーの限定も要件に適する。

この選択肢が正解になる条件

既存のOAIを継続する移行猶予が最優先で、オブジェクトがSSE-S3で暗号化され、KMS連携を必要としない非公開S3オリジンなら成立し得る。

C 正解

OAC、S3バケットポリシー、KMSキーポリシーを一貫して特定ディストリビューションに結び付け、非公開化、暗号化、HTTPS、キャッシュ、低運用負荷をすべて満たす。

D 不正解

S3を非公開にはできるが、配信がCloudFrontキャッシュを迂回し、URL発行処理と有効期限管理をアプリケーションへ追加するため、性能と運用負荷の条件に反する。

この選択肢が正解になる条件

利用者ごとに短時間だけ単一オブジェクトの直接ダウンロードを許可し、CloudFrontキャッシュの利用やURL発行負荷が問題にならない場合は適する。

7 今後使える判断基準

非公開S3をCloudFrontだけに公開する場合はOACとS3 RESTエンドポイントを第一候補にし、SSE-KMSならバケットポリシーだけでなくKMSキーポリシーまで認可経路を確認する。

8 関連サービスとの比較

OACはSigV4署名とSSE-KMSを含む現在のCloudFront向けオリジン制御に適する。OAIは従来構成の互換用途、presigned URLは利用者単位の一時的なS3直接アクセス、S3ウェブサイトエンドポイントは公開ウェブ配信向けである。

QUESTION 036

問題36

RESILIENT

複合難問

複数選択 (2つ)

Windows共有ファイルのMulti-AZ移行

製造会社はオンプレミスのWindowsファイルサーバークラスターで12TBの設計図を共有している。600人のWindowsクライアントは既存のself-managed Microsoft Active Directory (AD) で認証され、SMBとSIDに基づくNTFS ACLを使用し、\\files.example.local\designというパスを業務アプリに固定している。AWSから既存ADドメインコントローラーへの冗長なネットワーク到達性と必要なDNSは準備できる。AWS移行後は1つのAZの障害や計画メンテナンスでも共有を継続し、書き込み済みデータを失わず、通常のフェイルオーバーを30秒程度のI/O中断に抑えたい。アプリのパス変更や利用者・ACLの作り直しは認められず、独自のWindowsクラスターやレプリケーションは運用したくない。初回コピー中も元サーバーは更新されるため、停止時間は最終差分同期の2時間以内にする。移行先リージョンには2つ以上のAZがあり、DNSの手動切替やデータ競合の解消を運用員に依存させないことも条件である。要件を満たすために選択すべき対応を2つ選択してください。

- A** 既存のself-managed Microsoft ADへ参加させたFSx for Windows File ServerのMulti-AZファイルシステムを作成し、優先サブネットとスタンバイ側サブネットを別AZに配置する。
- B** 2つのAZにEFSマウントターゲットを作成し、WindowsクライアントにはNFSクライアントを追加する。NTFS ACLはアプリケーション側のユーザー表で再現する。
- C** 各AZにFSx for Windows File ServerのSingle-AZファイルシステムを1つずつ作成し、DFS Replicationを自社で構成して両方を同時書き込み可能にする。
- D** DataSyncでSMBソースからFSxへ初回と定期差分を転送し、切替時に最終差分を同期する。既存DNS名をFSxのDNSエイリアスとして関連付け、必要なSPNを構成する。

ここで答えを決める

正解・解説は次のページから

ANSWER 036

問題36の正解・解説

1 正解 **A・D**

- A** 既存のself-managed Microsoft ADへ参加させたFSx for Windows File ServerのMulti-AZファイルシステムを作成し、優先サブネットとスタンバイ側サブネットを別AZに配置する。
- D** DataSyncでSMBソースからFSxへ初回と定期差分を転送し、切替時に最終差分を同期する。既存DNS名をFSxのDNSエイリアスとして関連付け、必要なSPNを構成する。

2 問題文の重要な条件

- 既存ADの利用者・SID、SMB、NTFS ACL、既存UNCパスを維持する
- AZ障害と計画メンテナンスに対して管理された自動フェイルオーバーが必要である
- 更新中の12TBを反復同期し、最終停止時間を2時間以内に抑える

3 正解へ至る判断過程

1. WindowsネイティブのSMBとNTFSセマンティクスが必須なので、EFSではなくFSx for Windows File Serverが適合する。
2. 自前のレプリケーションを避けてAZ障害に耐えるには、スタンバイファイルサーバーを別AZに持つMulti-AZ構成を選ぶ。
3. 反復差分転送とパス不変の両方が必要なため、DataSyncによる段階移行とFSx DNSエイリアス／Kerberos SPNの構成を組み合わせる。

4 正解が要件を満たす理由

Aにより、FSxが別AZのファイルサーバーと同期ストレージを管理し、障害や保守時の自動フェイルオーバーを提供する。既存のself-managed ADへ参加するため、利用者SIDを再作成せずSMBとNTFS ACLを維持できる。Dにより、DataSyncで初回コピー後の変更だけを反復転送し、切替時の停止を最終差分に限定できる。DNSエイリアスとSPNを適切に構成すれば、固定UNC名とKerberos認証も維持できる。

5・6 各選択肢の検討

不適切な理由と、正解になり得る条件を対にして確認します。

A 正解

管理されたWindowsファイル共有のまま別AZへ自動フェイルオーバーし、既存ADの利用者SID、SMB、ACLという互換条件を満たす基盤である。

B 不正解

EFS RegionalはAZ耐障害性を持つがNFSファイルシステムであり、既存のSMB、NTFS ACL、WindowsのUNCパスをそのまま維持できない。

この選択肢が正解になる条件

クライアントがLinux中心でNFSとPOSIX権限を使用し、Windows固有のACLやSMB互換性が不要ならEFS Regionalが有力になる。

C 不正解

Single-AZを2台並べてもFSxの管理されたMulti-AZフェイルオーバーにはならず、DFS-Rの競合、監視、切替を自社運用するため要件に反する。

この選択肢が正解になる条件

各拠点が独立した名前空間を持ち、非同期複製のRPOとDFS-Rの運用を受け入れ、地域別の書き込み先を明確に分離できる場合は検討できる。

D 正解

オンライン中の初回・増分・最終同期で停止時間を短縮し、DNSエイリアスとSPNで固定パスとKerberos利用を保てる。

7 今後使える判断基準

Windows共有移行では、プロトコルと権限セマンティクス、障害単位、名前解決、差分移行を別々に確認する。SMB/NTFSとAZ自動フェイルオーバーならFSx for Windows Multi-AZ、短時間切替ならDataSync反復同期を組み合わせる。

8 関連サービスとの比較

FSx for WindowsはSMB、NTFS ACL、AD連携に適し、Multi-AZは管理された同期系フェイルオーバーを提供する。EFS RegionalはLinux/NFSの共有、2つのSingle-AZとDFS-Rは自己管理型の非同期複製、DataSyncは移行データ面を担う。

QUESTION 037

問題37

PERFORMANCE

本番相当

単一選択

Direct Connect経由の反復NFS転送

映像制作会社はオンプレミスNFSにある80TBの素材を、3か月かけてap-northeast-1のS3へ段階移行する。毎日約600GBが追加・更新され、初回転送後も夜間に差分だけを同期する必要がある。既存の10Gbps Direct ConnectとプライベートVIFを使用し、インターネット経由の転送は禁止されている。Direct Connectから到達可能なVPCにはDataSync用のInterface VPC Endpointを作成できる。ファイルの更新時刻や所有者など対応するメタデータを可能な範囲で維持し、転送後はチェックサムで整合性を検証したい。日中は制作トラフィックを優先するため転送帯域を制限し、失敗分は再実行できるようにする。夜間と日中で転送上限を変え、転送状況をCloudWatchから監視できることも求められる。専用の転送ソフトを開発せず、継続運用を最小化する構成はどれか。

- A** オンプレミスにS3 File Gatewayを配置し、既存NFSをゲートウェイの共有へ再マウントする。アプリケーションが全素材を読み書きすることでS3へ移行し、差分判定は利用者に任せる。
- B** aws s3 syncを実行するEC2インスタンスをAWS側に置き、NFSをDirect Connect越しにマウントする。cron、帯域制御、チェックサム台帳、失敗再開処理を自作する。
- C** AWS DMSのレプリケーションインスタンスからNFSをソースエンドポイントとして登録し、継続的レプリケーションでファイルをS3オブジェクトへ変換する。
- D** オンプレミスにDataSyncエージェントを配置し、Direct Connectから到達可能なVPCのDataSync Interface VPC Endpointをプライベートサービスエンドポイントとして使用する。NFSからS3へのタスクでスケジュール、帯域制限、増分転送、検証、再試行を管理する。

ここで答えを決める

正解・解説は次のページから

ANSWER 037

問題37の正解・解説

1 正解 **D**

- D** オンプレミスにDataSyncエージェントを配置し、Direct Connectから到達可能なVPCのDataSync Interface VPC Endpointをプライベートサービスエンドポイントとして使用する。NFSからS3へのタスクでスケジュール、帯域制限、増分転送、検証、再試行を管理する。

2 問題文の重要な条件

- 80TBの初回転送後、毎日更新される差分を反復転送する
- 10Gbps Direct ConnectとDataSyncのInterface VPC Endpointを使用し、インターネット経由を禁止する
- 帯域制限、整合性検証、失敗再実行を低い運用負荷で実現する

3 正解へ至る判断過程

1. 対象はデータベースではなくNFS上のファイルなので、NFSをネイティブなソースとして扱うオンライン転送サービスに絞る。
2. 初回以後は変更検出、帯域制御、検証、スケジュールが必要であり、単純なマウントや同期スクリプトよりDataSyncのタスク機能が適する。
3. DataSyncエージェントのアクティベーションとデータ転送をInterface VPC Endpointへ向け、プライベートVIFから到達させれば、公開サービスエンドポイントを経由せず管理機能を利用できる。

4 正解が要件を満たす理由

DはオンプレミスNFSを正式なソースとして扱い、初回フル転送後に変更ファイルだけを転送できる。タスク単位で帯域、スケジュール、ファイル処理、検証を設定し、実行結果を監視して再実行できるため、自作処理を減らせる。DataSyncのInterface VPC Endpointをプライベートサービスエンドポイントとして選び、プライベートVIFから到達させることで、エージェントの通信をインターネットへ出さずにDirect Connect上で完結できる。

5・6 各選択肢の検討

不適切な理由と、正解になり得る条件を対にして確認します。

A 不正解

S3 File GatewayはS3をNFS/SMB共有としてハイブリッド利用するサービスであり、既存NFSツリーを反復スキャンして移行・検証する専用機能ではない。再マウントも業務変更を伴う。

この選択肢が正解になる条件

移行ではなく、今後もオンプレミスアプリからS3オブジェクトをファイル共有として継続利用し、ローカルキャッシュが必要なら適する。

B 不正解

転送自体は可能だが、NFSの遠隔マウント、差分判定、メタデータ処理、帯域制御、検証、再開を自社実装・運用するため、低運用負荷の条件を満たさない。

この選択肢が正解になる条件

データ量が小さく一度限りで、特殊な変換ロジックのために同期処理を完全に制御する必要があり、開発運用を受け入れる場合は成立する。

C 不正解

DMSは対応するデータベースや一部データストアの移行・CDC用で、汎用NFSファイル共有をソースエンドポイントにはできない。

この選択肢が正解になる条件

移行対象がOracleやPostgreSQLなど対応DBの表で、変更データキャプチャにより停止時間を短縮した場合はDMSが正解になり得る。

D 正解

NFSからS3への大容量オンライン転送、増分、帯域制御、検証を一体化し、Interface VPC Endpoint経由でDirect Connectをプライベートに活用できる。

7 今後使える判断基準

ファイル移行で『オンライン、反復差分、検証、帯域制御』が並んだらDataSyncを優先する。Direct Connectのprivate VIFだけでDataSyncの公開エンドポイントへは到達できないため、インターネット禁止ならDataSync private service endpoint用のInterface VPC Endpointまで設計する。

8 関連サービスとの比較

DataSyncはNFS/SMBとAWSストレージ間の管理された転送、Storage GatewayはオンプレミスからAWSストレージを日常利用するハイブリッドアクセス、DMSはDB移行、CLI同期は小規模・特殊処理向けである。

QUESTION 038

問題38

COST

本番より難しい

単一選択

AZ間NAT転送費の削減

ECサイトは3つのAZに同数のEC2を配置し、各プライベートサブネットから外部決済API、S3、DynamoDBへ通信している。現在はAZ-aのパブリックサブネットにある1つのゾーン型NAT Gatewayへ全AZのデフォルトルートを向けている。月間30TBの送信のうち18TBはAZ-bとAZ-cから発生し、NATの処理料金に加えてAZ間データ転送料が増えている。1AZ障害時にも残るAZから外部APIへ接続できる必要がある。外部APIは固定の許可先が多く、NATインスタンスのパッチやフェイルオーバーは運用したくない。S3とDynamoDBへの通信はインターネットを経由させる必要がない。料金明細ではNAT処理量とRegional-DataTransferが主要因であり、外部APIへの実効通信量は今後も同程度で推移すると見込まれている。可用性を落とさず総コストを最も適切に削減する構成はどれか。

- A** VPCに自動モードのRegional NAT Gatewayを作成し、各AZのプライベートサブネットから同じNAT IDへルーティングする。S3とDynamoDBにはGateway VPC Endpointを追加する。
- B** NAT GatewayはAZ-aの1台だけを維持し、AZ-bとAZ-cのルートに同じNATを設定する。S3のライフサイクルで古いオブジェクトを低料金クラスへ移す。
- C** AZ-aのNAT Gatewayを1台のNATインスタンスに置き換え、Auto Recoveryを設定する。全AZのデフォルトルートはそのインスタンスへ向け、最大インスタンスサイズを選ぶ。
- D** 専用egress VPCにNAT Gatewayを1つ作成し、3AZのアプリVPCをTransit Gatewayで接続する。すべてのS3、DynamoDB、外部API通信を集中NATへ送る。

ここで答えを決める

正解・解説は次のページから

ANSWER 038

問題38の正解・解説

1 正解 **A**

- A** VPCに自動モードのRegional NAT Gatewayを作成し、各AZのプライベートサブネットから同じNAT IDへルーティングする。S3とDynamoDBにはGateway VPC Endpointを追加する。

2 問題文の重要な条件

- 高トラフィックが複数AZから発生し、現在は別AZのゾーン型NATを通過している
- AZ障害後も残存AZから外部APIへ接続できる必要がある
- NATの自己運用を避け、S3とDynamoDB通信のNAT処理料金も削減する

3 正解へ至る判断過程

1. 単一のゾーン型NATを別AZから使うと、NAT処理料金に加えてAZをまたぐ経路と単一AZ依存が生じる。
2. Regional NAT Gatewayはワークロードが存在するAZへ自動展開してゾーンアフィニティを保つため、単一のNAT IDで不要なAZ間転送と単一AZ依存を解消できる。
3. S3とDynamoDBはGateway VPC Endpointへ分離すれば、NATを通るデータ量をさらに減らし、プライベート経路を利用できる。

4 正解が要件を満たす理由

AのRegional NAT Gatewayは、ワークロードのある各AZへ自動的に拡張し、同一AZ内でNAT処理するゾーンアフィニティと高可用性を単一のNAT IDで提供する。これにより外部API通信の不要なAZ間転送と、AZ-aのゾーン型NATへの障害依存を解消できる。S3とDynamoDBをGateway VPC Endpointへ振り分ければ、その通信に対するNATデータ処理料金も不要になる。

5・6 各選択肢の検討

不適切な理由と、正解になり得る条件を対にして確認します。

A 正解

Regional NATの自動AZ展開とゾーンアフィニティでAZ間費用・障害依存を除き、Gateway EndpointでNAT処理量も削減する。

B 不正解

S3の保存料金は下がり得るが、問題の中心であるAZ間転送、NAT処理料金、AZ-a障害時の外部接続喪失を解消しない。

この選択肢が正解になる条件

通信がほぼAZ-a内だけで、NATのAZ障害を許容し、主な課題が長期保存S3オブジェクトのストレージ料金なら適切になり得る。

C 不正解

NATインスタンスはサイズ、パッチ、帯域、送信元/送信先チェック、障害時ルート切替を管理する必要があり、単一AZ経由と運用負荷の条件に反する。

この選択肢が正解になる条件

通信量が非常に少なくNAT Gatewayの時間料金が支配的で、独自機能が必要かつインスタンスのHA運用を受け入れる小規模環境なら候補になる。

D 不正解

集中egressは多数VPCの統制には有効だが、この単一VPCの高データ量ではTransit Gateway処理料金とAZ間経路を追加し、S3/DynamoDBまでNATへ送るため割高になる。

この選択肢が正解になる条件

多数のVPCで共通検査・許可リストを集中管理することが最優先で、追加のTransit Gateway処理料金を統制上の便益が上回る場合は有力である。

7 今後使える判断基準

複数AZのpublic NATでは、Regional NAT Gatewayを第一候補にして自動AZ展開・ゾーンアフィニティ・単一NAT IDの運用削減を評価する。private NATが必要ならゾーン型を使う。いずれもS3/DynamoDBはGateway Endpointへ先に逃がし、NAT処理量を減らす。

8 関連サービスとの比較

Regional NAT Gatewayはpublic NATをワークロードのあるAZへ自動展開し、単一IDで高可用性とゾーンアフィニティを提供する。ゾーン型NATはprivate NATやAZ単位の明示制御に使えるが、AZごとの作成と経路管理が必要である。NATインスタンスは柔軟だが自己運用、集中egress/TGWは多数VPCの統制と引き換えに処理料金が増える。

QUESTION 039

問題39

SECURE

複合難問

単一選択

暗号化SQSのクロスアカウント認可

金融グループでは、アカウントPの決済サービスがアカウントQの標準SQSキューへ取引イベントを送信し、アカウントCの不正検知ワーカーが受信・削除する。監査要件によりメッセージはアカウントQが管理するカスターマネージドKMSキーで暗号化し、Pの送信ロールとCの受信ロールだけに権限を限定しなければならない。各ロールのIAMポリシーには必要なSQSおよびKMSアクションを付与できる。現在、キュー所有者のIAMポリシーだけを変更してもPからSendMessageがAccessDeniedとなり、将来ほかの組織外アカウントへ権限が波及する構成も避けたい。Pは送信だけ、Cは受信と削除だけを行い、キュー属性変更やKMSキー管理を許可しないという職務分離も監査対象である。最小権限で正しく動作する設定はどれか。

- A** PとCの各ロールにSQSとKMSのIAM許可を付け、キュー側はアカウントQのデフォルト設定のままにする。アイデンティティポリシーだけでクロスアカウントを成立させる。
- B** キューポリシーでPの送信ロールにSendMessage、Cの受信ロールにReceiveMessage/DeleteMessage等を許可する。QのカスターマネージドKMSキーポリシーでも両ロールの必要な暗号操作を許可し、各ロールのIAM許可と組み合わせる。
- C** KMSキーポリシーでアカウントPとCのrootプリンシパルにkms:*を許可する。SQSキューポリシーは作成せず、KMSによる暗号化の認可がSQS APIの認可も兼ねるようにする。
- D** キューをAWSマネージドキーaws/sqsで暗号化し、SQS AddPermissionでアカウントPとCに全SQSアクションを付与する。キーのアクセス制御はAWSに委任する。

ここで答えを決める

正解・解説は次のページから

ANSWER 039

問題39の正解・解説

1 正解 **B**

- B** キューポリシーでPの送信ロールにSendMessage、Cの受信ロールにReceiveMessage/DeleteMessage等を許可する。QのカスタマーマネージドKMSキーポリシーでも両ロールの必要な暗号操作を許可し、各ロールのIAM許可と組み合わせる。

2 問題文の重要な条件

- SQSキュー、送信者、受信者が3つの異なるAWSアカウントにある
- アカウントQ所有のカスタマーマネージドKMSキーで暗号化する
- 特定の送受信ロールだけにSQS操作と暗号操作を最小権限で許可する

3 正解へ至る判断過程

1. クロスアカウントSQSでは、呼出側IAM許可だけでなく、キュー所有側のリソースベースポリシーが外部プリンシパルを許可する必要がある。
2. SSE-KMSではSQS APIの認可とKMSキー利用の認可は独立しているため、キューポリシーだけでもキーポリシーだけでも不足する。
3. カスタマーマネージドキーなら外部ロールを明示でき、外部アカウント側IAMポリシーと合わせて許可範囲を限定できる。

4 正解が要件を満たす理由

Bは、キューポリシーでデータプレーン操作を役割別に限定し、KMSキーポリシーで同じ外部ロールに必要な暗号操作だけを認める。さらにP/C側のIAMポリシーが実際の委任を制御するため、リソース所有側と呼出側の双方の認可を満たす。送信者に受信権限、受信者に送信権限を与えず、アカウント全体への広い許可も避けられる。

5・6 各選択肢の検討

不適切な理由と、正解になり得る条件を対にして確認します。

A 不正解

外部アカウントのIAMポリシーは、Qが所有するキューやKMSキーの信頼を単独では作れない。キューポリシーとKMSキーポリシーの許可が不足する。

この選択肢が正解になる条件

送信・受信ロール、キュー、KMSキーがすべて同一アカウントにあり、各リソースポリシーがIAMへの委任を妨げていない場合はIAM中心で設計できる。

B 正解

SQSリソースポリシー、KMSキーポリシー、外部IAMポリシーという三つの認可面を、特定ロールと必要操作に絞って満たす。

C 不正解

KMSの許可はSQS APIの許可を代替しない。さらに外部アカウントrootへのkms:*は過剰で、特定ロール限定という監査要件にも反する。

この選択肢が正解になる条件

暗号鍵管理を外部アカウントへ広く委任する必要がある、そのアカウント管理者がIAMで再委任する設計でも、別途SQSキューポリシーは必要になる。

D 不正解

AWSマネージドKMSキーのキーポリシーは顧客が編集できず、Q管理キーという要件を満たさない。アカウント単位的全SQSアクション許可も最小権限ではない。

この選択肢が正解になる条件

同一アカウント内の単純な暗号化キューで、顧客によるキーポリシー管理やロール単位のクロスアカウント制御が不要ならaws/sqsを利用できる。

7 今後使える判断基準

暗号化リソースのクロスアカウント問題は、サービスのリソースポリシー、KMSキーポリシー、呼出側IAMの三層に分解し、どれか一つで代替できると考えない。

8 関連サービスとの比較

SQSキューポリシーはSend/Receiveなどキュー操作を、KMSキーポリシーは暗号鍵利用を制御する。AWSマネージドキーは簡便だがポリシーを編集できず、外部主体を細かく許可する用途ではカスタマーマネージドキーが適する。

QUESTION 040

問題40

RESILIENT

本番相当

複数選択 (2つ)

起動が遅いEC2の事前初期化

保険会社の見積APIはALB配下のEC2 Auto Scalingグループで稼働している。インスタンスは起動後に暗号化済み基盤ルールセット20GBを取得して展開し、JITコンパイルと自己診断を行うため、リクエストを安全に処理できるまで14分かかる。展開済み基盤とJIT生成物は暗号化EBSへ永続化され、OS再起動後も再利用できる。通常は4台だがテレビ広告直後の5分間で12台必要になり、予測時刻は毎回ずれる。現在はスケールアウトが間に合わず、初期化途中のインスタンスが登録されて5xxも発生する。20GBの基盤はリリース期間中不変だが、1時間ごとに100MB未満の署名付き差分が公開される。事前初期化済みインスタンスなら、再開後に差分取得と検証を1分以内で完了できる。待機容量の費用は実行中より抑えつつ、初回リクエスト前に最新差分まで適用済みであることを保証したい。選択すべき対策を2つ選択してください。

- A** ヘルスチェック猶予期間を20分、単純スケーリングのクールダウンを20分にする。新規インスタンスは通常どおり起動し、必要台数に達するまで追加起動を止める。
- B** Auto Scalingグループに停止状態を使うwarm poolを追加し、ピーク前ではなく継続的に事前初期化済みインスタンスを確保して、スケールアウト時に再開する。
- C** 起動時とwarm poolからの復帰時にlifecycle hookでインスタンスを待機させ、必要な基盤初期化または最新差分の取得・検証が成功してからCompleteLifecycleActionを呼び出してInServiceへ進める。
- D** ALBのderegistration delayを14分に設定し、スケールアウト直後の新規インスタンスへ既存接続を長く残すことで初期化時間を吸収する。

ここで答えを決める

正解・解説は次のページから

ANSWER 040

問題40の正解・解説

1 正解 B・C

- B** Auto Scalingグループに停止状態を使うwarm poolを追加し、ピーク前ではなく継続的に事前初期化済みインスタンスを確保して、スケールアウト時に再開する。
- C** 起動時とwarm poolからの復帰時にlifecycle hookでインスタンスを待機させ、必要な基盤初期化または最新差分の取得・検証が成功してからCompleteLifecycleActionを呼び出してInServiceへ進める。

2 問題文の重要な条件

- EC2の初期化に14分かかる一方、需要増加は5分以内で予測時刻が不確実である
- 待機費用を実行中インスタンスより下げる必要がある
- 停止プールからの復帰後、1分以内の最新差分適用が成功するまでInServiceにしてはならない

3 正解へ至る判断過程

1. 需要を検知してから通常起動するだけでは14分を短縮できないため、事前に初期化した容量をグループ外の待機プールに用意する。
2. 停止状態のwarm poolはコンピュート実行料金を抑えながら、再開によって通常起動より速く容量を供給できる。
3. 停止中は毎時差分を取得できないため、復帰時のlifecycle hookで最新差分を適用し、成功イベントまでInService遷移を明示的に保留する。

4 正解が要件を満たす理由

Bにより、20GB基盤の展開、JIT、自己診断という長い初期化を済ませ、再利用物を暗号化EBSへ永続化したEC2を停止状態で保持できる。停止ではRAMは保持されないが、本問のJIT生成物はEBSから再利用できるため、新規起動より短時間で復帰する。停止中に公開された差分はCの復帰時lifecycle hookで1分以内に取得・検証し、成功通知までInServiceへの遷移を保留する。これにより5分の急増と最新性の両方を満たす。

5・6 各選択肢の検討

不適切な理由と、正解になり得る条件を対にして確認します。

A 不正解

猶予期間は早期の異常置換を遅らせるだけで起動を高速化せず、クールダウンは追加スケールを遅らせる可能性がある。処理準備完了の明示的保証にもならない。

この選択肢が正解になる条件

初期化が需要増加より十分短く、誤ったヘルスチェックによる置換やスケーリングの振動だけが問題なら猶予期間やクールダウン調整が有効である。

B 正解

EBSへ初期化成果物を永続化した停止インスタンスを用意し、不確実な急増へ通常起動より早く応答しながら待機中のコンピュータ費用を抑える。

C 正解

基盤初期化または復帰後の最新差分適用を確認するまで待機させ、古い・未準備のインスタンスがInServiceになることを防ぐ。

D 不正解

deregistration delayは登録解除される既存ターゲットの処理中接続をドレインする設定で、新規ターゲットの初期化や登録待ちには作用しない。

この選択肢が正解になる条件

スケールインやデプロイ時に旧インスタンスの長時間リクエストを切断せず完了させることが課題なら有効である。

7 今後使える判断基準

起動時間がスケール要求の猶予を超えるときは、warm poolで不変な重い初期化を前払いし、lifecycle hookで復帰時の差分適用を含む準備完了を保証する。停止中にも変わるデータが大きく、復帰後の追従が猶予内に終わらないなら停止warm poolだけでは解決しない。

8 関連サービスとの比較

warm poolは事前初期化容量の保持、lifecycle hookは状態遷移中のカスタム処理、health check grace periodは起動直後の誤置換防止、deregistration delayは終了側の接続ドレインを担う。

QUESTION 041

問題41

PERFORMANCE

本番より難しい

単一選択

仮想テープバックアップのクラウド移行

医療機関はオンプレミスのバックアップ製品から物理テープライブラリへ每晚18TBを書き込み、月次テープを7年間保管している。バックアップ製品はiSCSI接続のテープドライブとメディアチェンジャーを前提とし、今期中はソフトウェアもジョブ定義も変更できない。物理テープの搬送を廃止し、直近の仮想テープはローカルキャッシュから短時間で復元し、取り出した長期テープは低コストのアーカイブに置きたい。データセンターにはゲートウェイVM用の容量を用意でき、AWSへの回線断中はキャッシュ範囲でジョブを継続したい。バックアップ製品は一般的な仮想テープライブラリを認識でき、テープの取り出し操作を既存の保存ポリシーに組み込んでいる。既存インターフェイスを維持しつつ運用負荷を最も低くするサービスはどれか。

- A** S3 File Gatewayを配置し、バックアップ製品へNFS共有を提示する。各バックアップジョブをファイル出力へ変更し、S3 Lifecycleでアーカイブする。
- B** Cached Volume Gatewayを配置し、iSCSIブロックボリュームをバックアップ製品へ提示する。仮想テープとメディアチェンジャーはバックアップ製品内で自作する。
- C** Tape Gatewayを配置し、iSCSI-VTLとして仮想テープドライブとメディアチェンジャーを提示する。使用中テープをキャッシュし、取り出したテープをアーカイブする。
- D** Stored Volume Gatewayを配置し、全バックアップをオンプレミスiSCSIボリュームに保持する。EBSスナップショットを月次テープの代替として7年間保存する。

ここで答えを決める

正解・解説は次のページから

ANSWER 041

問題41の正解・解説

1 正解 C

- C** Tape Gatewayを配置し、iSCSI-VTLとして仮想テープドライブとメディアチェンジャーを提示する。使用中テープをキャッシュし、取り出したテープをアーカイブする。

2 問題文の重要な条件

- 既存バックアップ製品がiSCSIのテープドライブとメディアチェンジャーを要求する
- ジョブ定義を変更せず、物理テープ搬送を廃止する
- 直近データのローカル性能と、取り出したテープの低コスト長期保管を両立する

3 正解へ至る判断過程

1. 互換性の核心は単なるiSCSIブロックではなく、テープライブラリとして見えるVTLインターフェイスである。
2. Storage Gatewayのうち仮想テープドライブ、メディアチェンジャー、仮想テープを提示するのはTape Gatewayである。
3. ローカルキャッシュを使うアクティブテープと、取り出し後のクラウドアーカイブというライフサイクルもTape Gatewayの用途に一致する。

4 正解が要件を満たす理由

Cは既存ソフトウェアにiSCSI-VTLを提示するため、テープ向けジョブを変更せずに移行できる。書き込み中や頻繁に使うデータはゲートウェイのローカルキャッシュを利用し、仮想テープを取り出すとクラウド側のアーカイブへ移せる。物理媒体の管理と搬送をなくしながら、バックアップ製品が理解するテープ操作を維持する唯一の選択肢である。

5・6 各選択肢の検討

不適切な理由と、正解になり得る条件を対にして確認します。

A 不正解

File GatewayはNFS/SMBのファイル共有を提供するため、VTLを要求するバックアップ製品と互換ではなく、ジョブ定義変更禁止にも反する。

この選択肢が正解になる条件

バックアップ製品がNFS/SMB上のファイルをネイティブに出力でき、テープ操作を廃止してS3オブジェクトとして管理できる場合は適する。

B 不正解

Cached Volume GatewayはiSCSIでもディスクボリュームを提示する。テープドライブやメディアチェンジャーのセマンティクスを提供せず、自作は低運用負荷に反する。

この選択肢が正解になる条件

業務アプリがiSCSIディスクを必要とし、主要データをS3側へ置きつつ頻繁なブロックをローカルキャッシュしたい場合は最適になり得る。

C 正解

VTL互換、ローカルキャッシュ、仮想テープの取り出しと長期アーカイブを管理サービスとして提供し、既存ジョブを維持できる。

D 不正解

Stored Volumeはオンプレミスに全データを保持するiSCSIディスク用途であり、VTL互換性がない。18TB/日と7年分をローカル主媒体にする要件でもない。

この選択肢が正解になる条件

低遅延の全データをオンプレミスに保持するブロックアプリがあり、AWS側スナップショットを非同期バックアップとして使う場合に適する。

7 今後使える判断基準

Storage Gatewayはインターフェイスで選ぶ。NFS/SMBならFile、iSCSIディスクならVolume、既存バックアップソフトがテープドライブ／チェンジャーを要求するならTape Gatewayである。

8 関連サービスとの比較

Tape GatewayはVTL、File GatewayはS3またはFSxへのファイル共有、Cached Volumeはクラウド主データ+ローカルホットキャッシュ、Stored Volumeはオンプレミス主データ+クラウドスナップショットを提供する。

QUESTION 042

問題42

COST

複合難問

単一選択

多数VPCへの単一サービス公開

社内プラットフォーム部門は、プロバイダーVPCの3AZに配置したHTTPSのライセンスAPIを、30アカウント・120個のコンシューマーVPCへ提供する。買収した環境が含まれるためCIDRは重複している。利用者側からAPIの特定ポートへ接続できればよく、API側から利用者へ接続する必要はなく、コンシューマー同士の通信も禁止する。各VPCへ個別ルートを配布せず、新規VPCを月10個追加できる運用にしたい。インターネット公開は不可で、障害範囲をサービス単位に限定する。接続時間料金とデータ処理料金は許容するが、120個のフルメッシュ接続や各VPC全体を相互到達可能にするコストとリスクは避けたい。各コンシューマーは自身のエンドポイント利用料を負担でき、DNS名は共通化したいが、プロバイダー側のルート表へ利用者CIDRを登録してはならない。最適な構成はどれか。

- A** プロバイダーVPCと全コンシューマーVPCの間にVPC Peeringを作成し、双方の全CIDRをルートテーブルへ追加する。セキュリティグループでAPIポートだけを許可する。
- B** 全VPCを1台のTransit Gatewayへアタッチし、伝播を有効にして相互ルートを配布する。重複CIDRは静的ルートの優先順位で解決する。
- C** APIをインターネット向けALBで公開し、CloudFrontとAWS WAFを前段に置く。コンシューマーのNAT GatewayのElastic IPだけを許可する。
- D** APIの前にNLBを置いてVPC Endpoint Serviceとして公開し、許可したアカウントの各VPCにInterface VPC Endpointを作成する。Private DNSとエンドポイントのセキュリティ制御を用いる。

ここで答えを決める

正解・解説は次のページから

ANSWER 042

問題42の正解・解説

1 正解 **D**

- D** APIの前にNLBを置いてVPC Endpoint Serviceとして公開し、許可したアカウントの各VPCにInterface VPC Endpointを作成する。Private DNSとエンドポイントのセキュリティ制御を用いる。

2 問題文の重要な条件

- 120個のVPCに対して単一のHTTPSサービスだけを一方向に公開する
- コンシューマーVPCのCIDRが重複し、相互接続や推移ルーティングは不要である
- インターネット公開と大量のルート管理を避け、サービス単位で分離する

3 正解へ至る判断過程

1. VPC全体のネットワーク接続ではなく、特定サービスへの消費者開始接続だけが要件なので、ルーティング型接続よりサービス公開型接続が適する。
2. VPC PeeringとTransit Gatewayは重複CIDRを扱えず、接続数・ルート・到達範囲の管理も増える。
3. PrivateLinkのEndpoint ServiceとInterface Endpointはコンシューマー側のプライベートIPでサービスを提示し、CIDR重複や推移ルーティングを避けられる。

4 正解が要件を満たす理由

DはAWS PrivateLinkにより、NLB背後のAPIだけを各VPCのInterface Endpointへ公開する。コンシューマーからプロバイダーへのルート交換を必要とせず、重複CIDRでも利用できる。許可するアカウントとエンドポイント接続を制御し、コンシューマー間通信やプロバイダーからの逆方向接続を作らないため、サービス単位の最小到達性になる。接続数に応じた料金は生じるが、120個のネットワーク接続とルート運用を避ける要件に合う。

5・6 各選択肢の検討

不適切な理由と、正解になり得る条件を対にして確認します。

A 不正解

VPC Peeringは非推移で小規模な1対1接続には安価だが、重複CIDRをサポートせず、120本の接続とルート管理が必要になる。

この選択肢が正解になる条件

接続先が2〜3個の非重複CIDR VPCで、双方向の広いIP通信が必要かつルート管理が小規模なら低コストな選択になり得る。

B 不正解

Transit Gatewayも重複CIDR間のルーティングを解決せず、全VPCのアタッチ料金とデータ処理、ルート分離設計が必要で、単一API公開には範囲が広い。

この選択肢が正解になる条件

多数の非重複CIDR VPCとオンプレミスをハブ接続し、複数ポート・複数宛先の推移ルーティングや集中検査が必要なら適する。

C 不正解

WAFと送信元制限で保護は強化できるが、APIと通信経路をインターネット公開しないという明示条件に反し、各VPCのNAT費用も生じる。

この選択肢が正解になる条件

世界中の利用者へ公開するHTTP APIで、エッジ保護・キャッシュ・IP制限を優先し、公開エンドポイントが許容される場合は有効である。

D 正解

重複CIDRを許容しながら特定サービスだけを非公開で多数VPCへ提供し、ルート交換とVPC間の広い到達性を不要にする。

7 今後使える判断基準

接続対象が『ネットワーク全体』か『一つのサービス』かを最初に分ける。多数・重複CIDR・一方向・特定サービスならPrivateLink、ハブ型の広いIP接続ならTransit Gateway、少数VPCならPeeringを検討する。

8 関連サービスとの比較

PrivateLinkはサービス単位・重複CIDR対応・非推移、VPC Peeringは安価な少数1対1・非推移・非重複、Transit Gatewayは多数VPCの推移ルーティング・非重複、公開ALB/CloudFrontはインターネット利用者向けである。

QUESTION 043

問題43

SECURE

本番相当

単一選択

セキュリティ検出サービスの役割分担

小売会社はOrganizations配下の40アカウントについて、異なる種類のセキュリティリスクを継続検出したい。具体的には、S3に保存されたクレジットカード番号の所在、EC2とコンテナイメージに含まれる既知ソフトウェア脆弱性、盗まれたアクセスキーによる異常API操作や暗号通貨採掘、S3公開設定や暗号化設定が社内基準から逸脱した履歴を区別して把握する。単一製品ですべてを検出できるとは考えず、各サービスの結果は後でSecurity Hubに集約する。検出サービスはOrganizationsの委任管理者から有効化し、各アカウントのアプリ管理者には検出設定を停止させない方針である。不要な全データスキャンや誤ったサービス導入を避けるため、各一次検出サービスの役割を正しく対応付けた組み合わせはどれか。

- A** MacieでS3内の機密データを発見し、Inspectorでワークロードの脆弱性を評価し、GuardDutyで脅威・異常活動を検出し、AWS Configでリソース設定と準拠状態の変化を記録・評価する。
- B** GuardDutyでS3オブジェクト本文のカード番号を分類し、MacieでEC2パッケージのCVEを列挙し、Inspectorでアクセスキーの異常利用を検出し、Configでネットワークパケットを解析する。
- C** InspectorでS3の機密データと公開設定を同時に検出し、Configでマルウェアと暗号通貨採掘を検出する。MacieとGuardDutyは結果集約だけに使用する。
- D** Macieで全AWSリソースの設定逸脱を評価し、GuardDutyでOSパッチ適用状況を管理する。InspectorでCloudTrailを保管し、Configでカード番号の内容検査を行う。

ここで答えを決める

正解・解説は次のページから

ANSWER 043

問題43の正解・解説

1 正解 **A**

- A** MacieでS3内の機密データを発見し、Inspectorでワークロードの脆弱性を評価し、GuardDutyで脅威・異常活動を検出し、AWS Configでリソース設定と準拠状態の変化を記録・評価する。

2 問題文の重要な条件

- S3オブジェクト内容の機密データ発見が必要である
- ソフトウェア脆弱性、脅威・異常活動、設定準拠を別々に検出する
- Security Hubは集約先であり、各検出の一次サービスを正しく選ぶ

3 正解へ至る判断過程

1. データ内容、脆弱性、実行中の脅威、構成状態は観測対象が異なるため、一つのサービス名でまとめずテレメトリごとに対応付ける。
2. S3の機密データ分類はMacie、EC2やECR等の脆弱性評価はInspector、異常APIや悪意ある活動はGuardDutyが中心となる。
3. 設定変更の履歴とルールによる準拠評価はAWS Configであり、Security Hubはこれらの検出結果を関連付け・集約する。

4 正解が要件を満たす理由

Aは、各サービスが実際に観測する対象を正しく分離している。MacieはS3データを機械学習とパターン照合で分類し、Inspectorは対応ワークロードとイメージの脆弱性を評価する。GuardDutyはAWSの各種シグナルから悪意・異常を検出し、Configは設定項目の履歴とルール評価を提供する。結果をSecurity Hubへ集約すれば、一次検出の役割を混同せず組織全体を可視化できる。

5・6 各選択肢の検討

不適切な理由と、正解になり得る条件を対にして確認します。

A 正解

機密データ、CVE、脅威、構成準拠という四つの観測対象を、それぞれの専門サービスへ正しく割り当てている。

B 不正解

GuardDutyはカード番号の内容分類を主目的とせず、MacieはEC2のCVEスキャナーではない。InspectorとConfigの役割も逆である。

この選択肢が正解になる条件

名称をAの対応関係へ入れ替え、GuardDutyを異常活動、MacieをS3機密データ、Inspectorを脆弱性、Configを構成評価に割り当てた場合に正しくなる。

C 不正解

InspectorはS3本文分類を行わず、Configはマルウェアや採掘活動を挙動分析しない。MacieとGuardDutyは単なる集約サービスでもない。

この選択肢が正解になる条件

課題がEC2/ECR等の脆弱性だけならInspector、設定ルールの準拠だけならConfigを単独の中心にできるが、提示された四要件には追加サービスが必要である。

D 不正解

Macie、GuardDuty、Inspector、Configの対象をすべて取り違えており、どの検出要件も適切なテレメトリで満たせない。

この選択肢が正解になる条件

全リソースの構成逸脱をConfig、異常活動をGuardDuty、脆弱性をInspector、S3内容分類をMacieへ変更すれば成立する。

7 今後使える判断基準

セキュリティサービスは製品名ではなく『何を観測するか』で選ぶ。データ内容=Macie、ソフトウェア脆弱性=Inspector、脅威シグナル=GuardDuty、構成履歴・準拠=Config、結果集約=Security Hubと整理する。

8 関連サービスとの比較

MacieはS3機密データ、Inspectorはワークロード脆弱性、GuardDutyは脅威検出、AWS Configは構成評価、Security Hubは複数サービスの検出結果の集約・優先順位付けを担う。

QUESTION 044

問題44

RESILIENT

本番より難しい

複数選択 (2つ)

ALBを用いたimmutable blue/green

決済APIはALB配下のEC2 Auto Scalingグループで動作している。直近のインプレース更新ではパッケージの一部だけが更新され、インスタンスごとの差異と15分の停止が発生した。今後は各リリースを同一のゴールデンAMIから起動するimmutable方式とし、本番容量を維持したまま新バージョンを検証したい。ALBのホスト名は変えず、ヘルス指標と決済成功率を10分監視してから段階的に切り替え、異常なら数分以内に旧版へ戻す。デプロイ中の追加容量コストは許容するが、既存インスタンス上でファイルを更新してはならない。データベーススキーマは後方互換で、blueとgreenの同時実行が可能であり、DNS TTL変更による切替は採用しない方針である。要件を満たす設計要素を2つ選択してください。

- A** 既存Auto Scalingグループの各EC2へSystems Manager Run Commandで新パッケージを順次配布し、ALBのderegistration delay中に同じインスタンスを再起動する。
- B** 新AMIを使う別のAuto Scalingグループをgreen環境として起動し、別ターゲットグループへ登録する。テスト用リスナーまたは限定経路でgreenのヘルスと機能を検証する。
- C** 既存ターゲットグループからblueを全台同時に登録解除し、その同じグループへgreenを登録する。切替完了後にALBヘルスチェックを開始する。
- D** ALBの本番リスナーでblueとgreenのターゲットグループへの重みを段階的に変更し、CloudWatchアラーム等で失敗条件を監視する。異常時は重みをblueへ戻し、安定後にblueを終了する。

ここで答えを決める

正解・解説は次のページから

ANSWER 044

問題44の正解・解説

1 正解 **B・D**

- B** 新AMIを使う別のAuto Scalingグループをgreen環境として起動し、別ターゲットグループへ登録する。テスト用リスナーまたは限定経路でgreenのヘルスと機能を検証する。
- D** ALBの本番リスナーでblueとgreenのターゲットグループへの重みを段階的に変更し、CloudWatchアラーム等で失敗条件を監視する。異常時は重みをblueへ戻し、安定後にblueを終了する。

2 問題文の重要な条件

- 既存EC2を変更せず、新AMIから同一構成のgreen環境を作る
- 本番容量と同じALBホスト名を維持しながら事前検証と段階切替を行う
- メトリクス悪化時に旧blue環境へ数分以内で戻せる

3 正解へ至る判断過程

1. immutable要件では既存インスタンスを更新せず、新しいAMIから別フリートを作成する必要がある。
2. blueとgreenを別ターゲットグループにすれば、同時稼働、独立ヘルスチェック、事前検証が可能になる。
3. 本番リスナーの転送重みを変え、旧環境を一定時間残すことで段階移行と即時ロールバックを実現する。

4 正解が要件を満たす理由

Bは新AMIから独立したgreen Auto Scalingグループを作るため、既存個体の状態差を持ち込まないimmutableデプロイになる。別ターゲットグループにより本番切替前のヘルス確認もできる。DはALBの同じエンドポイントでトラフィック割合を小さく始めて増やし、メトリクスが悪化したらblue側へ戻せる。blueを検証期間中保持するため、容量低下や再構築待ちのないロールバックが可能である。

5・6 各選択肢の検討

不適切な理由と、正解になり得る条件を対にして確認します。

A 不正解

既存インスタンスを変更するインプレース／ローリング更新であり、個体差を排除するimmutable要件を満たさず、完全な旧環境も残らない。

この選択肢が正解になる条件

変更が小さく状態差を許容し、追加フリート費用を避けることが最優先で、ローリング中の容量低下を設計で吸収できる場合は選択肢になる。

B 正解

新AMI由来の別フリートと別ターゲットグループにより、immutable性、同時容量、事前検証を実現する。

C 不正解

blueを全台先に外すと停止や容量不足を招き、greenの事前ヘルス検証もできない。旧版へ戻すターゲット状態も失われやすい。

この選択肢が正解になる条件

完全停止を伴う一括切替が許容され、テスト環境でgreenを別途十分検証済みで、迅速なロールバック要件がない場合に限り簡素化案になり得る。

D 正解

二つのターゲットグループ間で段階的に本番トラフィックを移し、監視結果に応じて旧版へ即座に戻せる。

7 今後使える判断基準

immutable blue/greenでは『新しいフリート』『別ターゲットグループ』『段階的なトラフィック制御』『旧フリート保持』の四点を確認する。ローリング更新や接続ドレインだけでは代替できない。

8 関連サービスとの比較

blue/greenは追加容量と引き換えに事前検証・高速ロールバックを得る。immutable updateは新規個体へ置換して状態差を防ぎ、in-place/rollingは安価だが旧環境の即時復帰性と完全な同一性が弱い。

QUESTION 045

問題45

PERFORMANCE

複合難問

単一選択

回線制約下の大容量オフライン移行

研究機関は沿岸観測所のNASにある900TBの読み取り専用データを、6週間以内にS3へ移行する。観測所の回線は200Mbpsで日中の半分を観測通信に予約する必要がある、増速工事には9か月かかる。同機関は既存のSnowball Edge利用資格と受取体制を持つ。データはKMS管理下で暗号化し、輸送状況を追跡し、S3取り込み後に完全性を確認する。ファイルの30%は1MB未満で、基準時点以後に追加・変更される約2TBだけは最終切替時にオンライン転送できる。NASは変更ジャーナルを保持し、基準時点後に変わった全ファイルの相対パスをDataSync用CSV manifestへ出力できる。初回と差分を明確に分離し、現地へ常駐サーバーを増設せず、期限とセキュリティを満たす最適な方法はどれか。

- A** 既存200Mbps回線でDataSyncを24時間実行し、帯域制限を100Mbpsにする。900TBの初回コピーと最終差分を同じタスクで6週間以内に完了させる。
- B** 必要容量分のSnowball Edgeデバイスを並行利用し、元の相対パスとS3キーの対応を保って基準時点の全データをコピーする。返送・取込み後に検証し、変更ジャーナルから作ったmanifestをDataSyncへ指定して約2TBだけを送る。
- C** 観測所からS3 Transfer Accelerationのエンドポイントへマルチパートアップロードする。エッジロケーションを使えば物理回線容量を超えて900TBを転送できる。
- D** 1GbpsのDirect Connect専用接続を新設してDataSyncを実行する。工事完了を待ち、残り期間はS3 Lifecycleで転送速度を上げる。

ここで答えを決める

正解・解説は次のページから

ANSWER 045

問題45の正解・解説

1 正解 **B**

- B** 必要容量分のSnowball Edgeデバイスを並行利用し、元の相対パスとS3キーの対応を保って基準時点の全データをコピーする。返送・取込み後に検証し、変更ジャーナルから作ったmanifestをDataSyncへ指定して約2TBだけを送る。

2 問題文の重要な条件

- 900TBを6週間以内に移す一方、利用可能WAN帯域は実質約100Mbpsに限られる
- 既存のSnowball Edge利用資格、物理受取体制、KMS暗号化と追跡要件がある
- 変更ジャーナルから正確なmanifestを作り、初回後の2TBだけをDataSyncでオンライン転送できる

3 正解へ至る判断過程

1. 利用可能帯域で900TBを転送する理論時間は期限を大幅に超えるため、オンライン高速化だけでは解決できない。
2. 初回の静的な大容量を複数の物理デバイスへファイル範囲で分割して並列搬送し、元パスとS3キーの対応を維持する。
3. 物理取り込み後はDataSyncの通常の変更判定に初回Snowball転送を認識させるのではなく、変更ジャーナル由来のmanifestで約2TBを明示して送る。

4 正解が要件を満たす理由

BはWAN帯域を初回900TBのボトルネックから外し、複数デバイスへの並列コピーと物理輸送で期限内移行を現実的にする。Snowball Edgeは暗号化とジョブ追跡を提供し、返送後にS3へ取り込める。Snowball取込みにはDataSyncがNFSからS3へ付与する変更判定用メタデータがないため、DataSyncのCHANGED設定だけに初回分の識別を任せない。変更ジャーナルから作成したmanifestで基準時点後の約2TBを明示し、検証付きで転送する。

5・6 各選択肢の検討

不適切な理由と、正解になり得る条件を対にして確認します。

A 不正解

100Mbpsで900TBを送るには理論値でも6週間をはるかに超え、プロトコル効率や共有回線を考慮するとさらに長い。DataSyncも物理帯域を超えられない。

この選択肢が正解になる条件

データ量が数十TB以下で期限に十分余裕があり、継続的な差分をオンラインで送ることが主目的ならDataSync単独が適する。

B 正解

大容量初回分を安全な物理転送へ、少量差分をオンライン転送へ分け、帯域・期限・整合性・運用制約をすべて満たす。

C 不正解

Transfer Accelerationは長距離インターネット経路を最適化するが、観測所の200Mbpsというラストマイル容量を超えることはできない。

この選択肢が正解になる条件

十分なインターネット帯域があり、遠距離から中規模ファイルをS3へ高速アップロードする際の経路遅延が主なボトルネックなら有効である。

D 不正解

Direct Connectは安定した帯域を提供し得るが、工事に9か月かかるため6週間の期限を満たさない。Lifecycleは保存階層を変える機能で転送速度を上げない。

この選択肢が正解になる条件

恒常的な大容量通信が今後も続き、回線開通まで待てる長期計画であればDirect ConnectとDataSyncの組み合わせが適する。

7 今後使える判断基準

大容量移行はデータ量÷実効帯域で下限時間を計算し、期限を超えるなら初回を物理転送、切替差分をオンライン転送へ分ける。初回がDataSync以外ならDataSync固有メタデータを前提に差分判定せず、変更ジャーナル、スナップショット差分、manifestなどで送るファイルを明示する。

8 関連サービスとの比較

Snowball Edgeは既存利用者の回線制約下の物理大容量転送、DataSyncはオンラインの初回・差分転送、Transfer Accelerationは十分なインターネット帯域がある遠距離S3転送、Direct Connectは継続的な専用接続向けである。

QUESTION 046

問題46

COST

本番相当

単一選択

履歴のない変動負荷とDynamoDB課金モード

旅行会社は新しい座席照会APIの状態ストアとしてDynamoDBを採用する。通常は每秒100回未満だが、航空会社のセール開始時には予告なく数分で每秒20,000回まで増え、その後ほぼゼロになる。新サービスのため十分な履歴がなく、ピーク時のスロットリングは売上損失につながる。開発チームはキャパシティ予測とスケーリングポリシー調整を行う要員を持たず、使われていない時間のプロビジョンド容量にも支払いたくない。キー分布は負荷試験で均等であることを確認済みで、半年後に負荷が安定したら再評価する。ピークの正確な開始時刻や直前値は予測できず、数週間後には前回ピークを超える可能性もあるため、初期の低い設定値に依存したくない。現時点で最も適切な容量モードはどれか。

- A** Provisionedモードで平均100 RCU/WCUだけを設定し、スロットリングされた要求はアプリケーションが無制限に即時再試行する。Auto Scalingは使用しない。
- B** ProvisionedモードとApplication Auto Scalingを使用し、目標使用率を高く設定する。初期容量は平常値にし、急増後にメトリクスを見て段階的に増やす。
- C** On-Demandモードを使用し、負荷試験で求めた読取・書込内訳に合わせて販売前にWarmThroughputを20,000回/秒相当以上へ事前ウォームする。CloudWatchで消費量とスロットリングを監視し、安定後にProvisionedを比較する。
- D** Provisionedモードを平常値に固定し、DAXクラスターを追加してすべての読み取りと書き込みの急増を吸収する。DynamoDB側の容量は増やさない。

ここで答えを決める

正解・解説は次のページから

ANSWER 046

問題46の正解・解説

1 正解 C

- C** On-Demandモードを使用し、負荷試験で求めた読取・書込内訳に合わせて販売前にWarmThroughputを20,000回/秒相当以上へ事前ウォームする。CloudWatchで消費量とスロットリングを監視し、安定後にProvisionedを比較する。

2 問題文の重要な条件

- 新規ワークロードで履歴がなく、平常時と突発ピークの差が非常に大きい
- スロットリングを避ける一方、アイドル時のプロビジョンド容量料金を抑える
- キャパシティ計画とAuto Scaling調整の運用負荷を最小化する

3 正解へ至る判断過程

1. ホットパーティションではなくテーブル全体の予測不能な変動が問題なので、キー再設計より容量モードを評価する。
2. Provisioned Auto Scalingは継続的・比較的緩やかな変動には有効だが、メトリクスに反応するまでの急増を初期容量だけで保証しにくい。
3. On-Demandは実リクエスト課金でアイドル容量を抱えず、必要なWarmThroughputをピーク水準へ事前設定すれば、新規テーブルの既定初期スループットや直前ピークの2倍という即時拡張の目安にも依存しない。

4 正解が要件を満たす理由

Cは販売前に読取・書込別のWarmThroughputを必要ピークまで引き上げるため、履歴のない新規On-Demandテーブルでも20,000回/秒相当を即時に受ける前提を用意できる。その後は処理したリクエストへの課金となり、長いアイドル中にProvisioned容量を抱えない。事前ウォーム費用を含めても売上損失回避を優先でき、負荷安定後はProvisionedとの総額比較へ移れる。

5・6 各選択肢の検討

不適切な理由と、正解になり得る条件を対にして確認します。

A 不正解

平均値だけの固定容量では20,000回/秒のピークを処理できず、無制限の即時再試行はスロットリングと負荷をさらに悪化させる。

この選択肢が正解になる条件

要求量が常に100単位以内で、超過時の遅延や指数バックオフ付き再試行を許容し、固定容量が十分なら簡素なProvisioned構成にできる。

B 不正解

Auto Scalingはコスト効率を改善できるが、予告なく数分で200倍になる初動ではスケールが追いつかず、初期容量を低くすると売上要件を満たしにくい。

この選択肢が正解になる条件

日次パターンが予測でき、変化が緩やかで、十分な最小容量やスケジュールスケールリングを設定できる定常高負荷ならProvisioned+Auto Scalingが有利になる。

C 正解

必要ピークまで事前ウォームして未知の開始時刻に備えつつ、長いアイドル時間に未使用Provisioned容量を持たないという全条件に合う。

D 不正解

DAXは読み取りキャッシュであり書き込み容量を代替しない。クラスターの常時料金も発生し、照会結果のキャッシュ適合性も示されていない。

この選択肢が正解になる条件

DynamoDBへの反復的で結果整合性を許容する読み取りが支配的で、マイクロ秒レベルの応答が必要かつDAX常時費用を正当化できる場合は有効である。

7 今後使える判断基準

DynamoDBは、未知・予測不能・アイドルが長いならOn-Demand、予測可能・持続的・高利用率ならProvisioned+Auto Scalingを基準にする。ただしOn-Demandも無限に瞬時拡張するわけではない。新規テーブルの大きな初回ピークは、読取・書込別の必要量を算出してWarmThroughputを事前に引き上げる。

8 関連サービスとの比較

On-Demandはリクエスト課金で容量運用を減らせるが、既定の初期WarmThroughputと過去ピークに基づく即時拡張の目安がある。WarmThroughputの事前ウォームは大きな初回ピークに備える機能で追加費用を伴う。Provisioned+Auto Scalingは安定利用で単価を下げやすく、DAXは容量モードではなく読み取りキャッシュである。

QUESTION 047

問題47

SECURE

本番より難しい

単一選択

外部IdPと複数アカウント権限の集約

グローバル企業はAWS Organizations内の70アカウントを使用し、従業員8,000人の認証元をMicrosoft Entra IDに統一している。入退社と部署異動はIdPのグループで管理し、AWSには長期IAMユーザーを作成しない。開発者、監査者、運用管理者の職務ごとに同じ権限テンプレートを複数アカウントへ割り当て、1時間の一時セッションを発行したい。退職者はIdPで無効化した後、AWS側の70アカウントを個別修正せずアクセスを失う必要がある。管理アカウントの利用も同じ仕組みに含め、緊急用資格情報を除いて個別アカウントのサインインURLやロール一覧を利用者に管理させない。SAML連携とユーザー／グループの自動同期を利用し、権限変更の運用を最小化する最適な構成はどれか。

- A** 各AWSアカウントにSAML IdPと同名のIAMロールを個別作成し、70個の信頼ポリシーと権限ポリシーを各アカウント管理者が手作業で同期する。
- B** Entra IDの全利用者と同名のIAMユーザーを各AWSアカウントへ作成し、アクセスキーをSecrets Managerに保存する。部署異動時はIAMグループを更新する。
- C** Amazon Cognito user poolをEntra IDへフェデレーションし、identity poolから全従業員へ共通の管理者IAMロールを払い出してAWSコンソールへ接続する。
- D** OrganizationsのIAM Identity Centerを有効にし、Entra IDを外部IdPとしてSAML連携しSCIMでユーザーとグループを同期する。職務別permission setを作り、グループを対象アカウントへ割り当てる。

ここで答えを決める

正解・解説は次のページから

ANSWER 047

問題47の正解・解説

1 正解 **D**

- D** OrganizationsのIAM Identity Centerを有効にし、Entra IDを外部IdPとしてSAML連携しSCIMでユーザーとグループを同期する。職務別permission setを作り、グループを対象アカウントへ割り当てる。

2 問題文の重要な条件

- 70アカウントと8,000人を中央の外部IdPグループから管理する
- 長期IAMユーザーを作らず、職務別の一時セッションを発行する
- 入退社・異動と複数アカウントへの権限テンプレート配布を自動化する

3 正解へ至る判断過程

1. 対象はB2Cアプリ利用者ではなく、従業員による複数AWSアカウントへのワークフォースアクセスである。
2. IAM Identity CenterはOrganizationsのアカウント割当とpermission setを中央管理し、一時的なAWSアクセスを提供する。
3. SAMLは認証、SCIMはユーザー・グループのプロビジョニングを担うため、外部IdP側のライフサイクルをAWS側へ反映できる。

4 正解が要件を満たす理由

Dは、既存IdPを認証元として維持しつつ、SCIMで利用者とグループの追加・更新・削除をIdentity Centerへ同期する。permission setを職務単位で一度定義して複数アカウントへ割り当てられ、各アカウントには対応ロールが管理される。利用者は長期アクセスキーではなく一時セッションを得るため、退職や異動をIdPグループと割当の変更から中央反映できる。

5・6 各選択肢の検討

不適切な理由と、正解になり得る条件を対にして確認します。

A 不正解

SAMLフェデレーション自体は可能だが、70アカウントのIdP、ロール、ポリシーを個別管理し、職務テンプレートのドリフトを招くため運用要件を満たさない。

この選択肢が正解になる条件

対象が1〜2アカウントで、既存の直接SAMLロール連携を変更できず、中央アカウント割当機能が不要なら実用的な構成になり得る。

B 不正解

多数の長期IAMユーザーとアクセスキーを複製し、認証元とライフサイクルを二重化するため、明示された禁止条件と最小運用に反する。

この選択肢が正解になる条件

AWSサービス外のレガシーツールがどうしてもIAMユーザーの長期キーを要求し、短期資格情報へ移行できない限定的な例外なら、厳格なローテーション付きで検討される。

C 不正解

Cognitoは主に顧客向けアプリ認証とアプリからのAWS資格情報発行用であり、従業員の70アカウントへの職務別コンソール権限管理には適さない。共通管理者ロールも過剰である。

この選択肢が正解になる条件

対象がモバイル／Webアプリの顧客で、認証後にS3など限定AWSリソースへ一時資格情報を渡すB2C要件ならCognitoが適する。

D 正解

外部IdP、SCIM、permission set、アカウント割当を組み合わせ、認証・利用者同期・権限配布・一時セッションを中央化する。

7 今後使える判断基準

従業員の複数AWSアカウントアクセスはIAM Identity Center、外部IdPの認証はSAML、利用者・グループ同期はSCIM、権限テンプレートはpermission setと役割を分けて判断する。

8 関連サービスとの比較

IAM Identity Centerはワークフォースのマルチアカウントアクセス、直接IAM SAMLロールは少数アカウントの個別連携、Cognitoはアプリの顧客ID、IAMユーザーは例外的な長期主体向けである。

QUESTION 048

問題48

RESILIENT

複合難問

複数選択 (2つ)

FIFOの順序単位と高スループット

証券会社は口座ごとの注文イベントをSQSで処理する。ある口座の注文は発生順に必ず1つずつ処理する必要があるが、異なる口座は並列処理できる。ピークは毎秒6,000メッセージで、送信APIのタイムアウト再試行により同じ注文が再送されることがある。注文IDは全体で一意で、5分以内の再送をキューで抑止したい。コンシューマーは処理結果を注文IDで冪等に記録する。同じ口座に対する次の注文は前の注文が削除されるまで処理を開始してはならないが、口座をまたぐ全体順序には業務上の意味がない。単一のメッセージグループによるボトルネックを避けながら、FIFOの順序と重複排除を利用するために必要な設定を2つ選択してください。

- A** すべての注文に同じMessageGroupIdを設定し、コンシューマー数を増やす。キュー内の全注文を厳密な全体順序にすることで6,000件/秒を並列処理する。
- B** 標準SQSキューを使用し、SentTimestampでコンシューマー側ソートを行う。Visibility Timeout中に同じ注文が再配信されないことを重複排除の保証とする。
- C** FIFOキューで高スループット設定を有効にし、deduplication scopeをmessage group、FIFO throughput limitをper message group IDとして、多数のグループへ分散する。
- D** MessageGroupIdに口座IDを設定し、MessageDeduplicationIdに一意な注文IDを設定する。各口座内は直列、口座間は並列にし、コンシューマーの冪等性も維持する。

ここで答えを決める

正解・解説は次のページから

ANSWER 048

問題48の正解・解説

1 正解 C・D

- C** FIFOキューで高スループット設定を有効にし、deduplication scopeをmessage group、FIFO throughput limitをper message group IDとして、多数のグループへ分散する。
- D** MessageGroupIdに口座IDを設定し、MessageDeduplicationIdに一意的な注文IDを設定する。各口座内は直列、口座間は並列にし、コンシューマーの冪等性も維持する。

2 問題文の重要な条件

- 順序保証の範囲は全注文ではなく同一口座内である
- 毎秒6,000件を異なる口座間で並列処理する
- 5分以内の送信再試行を注文IDで重複排除し、処理側も冪等にする

3 正解へ至る判断過程

1. FIFOの並列性はMessageGroupId単位なので、業務上の順序境界である口座IDをグループにする。
2. 全口座を一つのグループへ入れると不要な全体直列化が起きるため、多数のグループへ均等に分散する。
3. 高スループットFIFOのグループ単位設定と、注文IDによる明示的重複排除を組み合わせ、キュー外でも冪等性を維持する。

4 正解が要件を満たす理由

Cは高スループットFIFOで複数のMessageGroupIdをパーティションへ分散し、グループ単位のスループットと重複排除を利用できるようにする。Dは順序が必要な業務キーである口座IDをMessageGroupIdにするため、同一口座だけを逐次処理し、別口座を並列化できる。注文IDをDeduplicationIdにすれば5分の重複排除ウィンドウ内の再送を抑え、時間外や処理再試行はコンシューマーの冪等性で補える。

5・6 各選択肢の検討

不適切な理由と、正解になり得る条件を対にして確認します。

A 不正解

一つのMessageGroupId内は順序維持のため同時処理されず、コンシューマーを増やしても全口座が直列化されスループットを活用できない。

この選択肢が正解になる条件

業務上すべてのメッセージに単一の厳密な全体順序が必要で、低いスループットを受け入れられる場合は単一グループが正しい。

B 不正解

標準キューは順序と重複なしを保証せず、タイムスタンプの事後ソートでは並行処理中の欠落・遅延レコードを安全に扱えない。Visibility Timeoutも重複排除機能ではない。

この選択肢が正解になる条件

順序が不要で少なくとも1回配信を前提に完全な冪等処理ができ、最大スループットと低コストを優先するなら標準キューが適する。

C 正解

高スループットFIFOをグループ単位で有効にし、多数の口座グループへ負荷を分散できる。

D 正解

口座を順序境界、注文IDを重複境界として正しくモデル化し、口座間並列性と処理冪等性を両立する。

7 今後使える判断基準

FIFO設計では、MessageGroupIdを『順序が必要な最小業務単位』、DeduplicationIdを『同じ送信と判定する一意キー』にする。高スループットが必要なら多数のグループへ分散し、5分を超える重複はアプリ冪等性で守る。

8 関連サービスとの比較

FIFOはグループ内順序と送信重複排除、標準SQSは高い並列性と少なくとも1回配信を提供する。Visibility Timeoutは処理中の再取得抑制、DLQは反復失敗の隔離であり、重複排除の代替ではない。

QUESTION 049

問題49

PERFORMANCE

本番相当

単一選択

NFSをEFSとS3へ分割するオンライン移行

大学はオンプレミスNFSの35TBをAWSへ移行する。うち8TBはLinux研究アプリが移行後もPOSIX権限と共有ファイル操作で利用するためEFSへ、残り27TBは更新されない成果物でアプリをS3 API対応に変更済みのためS3へ保存する。利用者は初回コピー中もNFSを更新し、切替停止は4時間以内にしたい。ファイル数は4,000万で、所有者、モード、更新時刻をEFS側で可能な限り保つ。転送の帯域制限、結果レポート、整合性検証が必要で、rsyncサーバーを長期運用したくない。ネットワーク帯域は初回を数日で送れるが、一度の長時間停止は許されず、移行後はオンプレミス装置を速やかに廃止する予定である。最も適切な移行方法はどれか。

- A** オンプレミスにDataSyncエージェントを配置し、NFSからEFS用とS3用に対象パスを分けたタスクを作る。初回、定期増分、読み取り専用化後の最終増分と検証を実行する。
- B** S3 File GatewayをNFSの前段に置き、全35TBを一度同じS3バケットへ書き込む。切替後は同じバケットをEFSとしてマウントしPOSIX権限を復元する。
- C** Snowball Edgeに全データを一度コピーし、S3取り込み完了後にS3から8TBを手作業でEFSへコピーする。初回コピー後の変更は移行対象外にする。
- D** AWS側EC2へNFSをマウントし、常時稼働のrsyncデーモンでEFSへ同期する。S3分もEFSへ置き、後日EC2スクリプトでオブジェクト化する。

ここで答えを決める

正解・解説は次のページから

ANSWER 049

問題49の正解・解説

1 正解 **A**

- A** オンプレミスにDataSyncエージェントを配置し、NFSからEFS用とS3用に対象パスを分けたタスクを作る。初回、定期増分、読み取り専用化後の最終増分と検証を実行する。

2 問題文の重要な条件

- 共有ファイル用途はEFS、変更されないオブジェクト用途はS3へ分ける
- 移行中も更新されるため、初回後の増分と短い最終切替が必要である
- 大量ファイルのメタデータ、帯域、検証、レポートを管理サービスで扱う

3 正解へ至る判断過程

1. 移行後のアクセスセマンティクスが異なるため、全データを一つの保存先へ無理に統一せずパスごとにEFSとS3を選ぶ。
2. DataSyncはNFSをソースにEFSとS3の両方を宛先として扱い、タスク再実行時に増分転送できる。
3. 初回と複数の増分を業務中に実行し、最後だけソースを読み取り専用にして差分を送ると停止時間を短縮できる。

4 正解が要件を満たす理由

Aはデータセットの用途別にタスクとフィルターを分け、EFSへは対応するファイルメタデータを保ちながら共有ファイルとして、S3へはオブジェクトとして直接転送できる。初回転送後に変更分だけを繰り返し同期し、最終メンテナンスではNFSを読み取り専用にして残差を転送・検証できる。帯域制御、タスクレポート、CloudWatch監視も利用でき、常設の自作rsync基盤を避けられる。

5・6 各選択肢の検討

不適切な理由と、正解になり得る条件を対にして確認します。

A 正解

NFSからEFS/S3への用途別直接転送、反復増分、メタデータ処理、検証を一つの管理サービスで実現する。

B 不正解

File Gatewayはハイブリッド共有を提供するが、S3バケットをEFSとしてマウントすることはできず、S3オブジェクトとPOSIX共有ファイルのセマンティクスも同一ではない。

この選択肢が正解になる条件

移行後もオンプレミスからS3をNFS/SMB共有として継続利用し、EFSのPOSIX共有が不要ならFile Gatewayが適する。

C 不正解

物理転送は初回大容量には使えるが、更新差分を除外するため4時間切替と完全性を満たさず、S3からEFSへの追加コピーも増やす。

この選択肢が正解になる条件

データが完全に静止しており回線では期限に間に合わず、S3だけが最終宛先である大容量一括移行なら適する。

D 不正解

rsyncの設計、スケール、再開、検証、監視を自己運用し、S3向けデータまで高価なEFSへ一時保管して二重転送するため運用・コストが増える。

この選択肢が正解になる条件

DataSyncが扱えない特殊な変換や除外ロジックが必須で、対象が小さく、自作転送基盤を保守する要員がいる場合は選択肢になる。

7 今後使える判断基準

移行先は現在の保存装置ではなく移行後のアクセス方式で分ける。NFS共有を保つデータはEFS、APIで扱う不変データはS3、オンライン反復移行はDataSyncで初回・増分・最終差分に分ける。

8 関連サービスとの比較

DataSyncは移行と反復同期、EFSはPOSIX/NFS共有、S3はオブジェクト保存、File GatewayはオンプレミスからS3をファイル共有として利用する橋渡し、Snowballは静的な大容量一括転送向けである。

QUESTION 050

問題50

COST

本番より難しい

単一選択

変動するAurora容量の細粒度スケーリング

チケット販売会社はAurora PostgreSQL互換の購入DBを運用する。平日は2～4 vCPU相当で十分だが、発売開始時刻は週ごとに異なり、数分で通常の12倍へ増えて2時間後に戻る。月の85%は低負荷である。接続を切断する長い一時停止は許容されず、Multi-AZのストレージ耐久性とSQL互換性を維持したい。担当者が発売ごとにDBクラスを変更する運用は避け、上限を設定して予算超過を防ぎたい。一方、年間を通じて最大サイズのインスタンスを常時2台持つ費用は削減したい。DB接続先エンドポイントとアプリのSQLは変更せず、容量変更中にもトランザクションを継続できることを負荷試験で確認する方針である。最も適切な構成はどれか。

- A** 最大ピークに合わせたプロビジョンドAurora writerと同サイズreaderを常時稼働し、DBクラス変更は年1回だけ行う。未使用容量は可用性の費用として受け入れる。
- B** Aurora Serverless v2のDBインスタンスを使用し、通常処理を維持できる最小ACUと予算内の最大ACUを設定する。必要に応じてreaderも配置し、CloudWatchでACUと接続を監視する。
- C** プロビジョンドAuroraのwriterを平常時の最小DBクラスにし、EventBridge SchedulerとLambdaで毎週同じ時刻に最大クラスへ変更して、発売後に元へ戻す。クラス変更中の再接続は許容する。
- D** 標準RDS PostgreSQLのSingle-AZインスタンスへ変更し、EC2 Auto Scalingで同じEBSボリュームを複数DBインスタンスからマウントして負荷を分散する。

ここで答えを決める

正解・解説は次のページから

ANSWER 050

問題50の正解・解説

1 正解 **B**

- B** Aurora Serverless v2のDBインスタンスを使用し、通常処理を維持できる最小ACUと予算内の最大ACUを設定する。必要に応じてreaderも配置し、CloudWatchでACUと接続を監視する。

2 問題文の重要な条件

- 負荷が予測しにくく短時間で大きく変動し、低負荷時間が85%を占める
- SQL互換性、接続継続、Auroraの可用性を維持する
- 手動リサイズとピーク容量の常時費用を避け、最大コストを制御する

3 正解へ至る判断過程

1. ピークが持続せず開始時刻も一定でないため、最大プロビジョンド容量の常時保持は利用率が低い。
2. Aurora Serverless v2は同じAuroraアーキテクチャ内でACUを細かく増減し、プロビジョンドDBクラスの置換より負荷へ追従しやすい。
3. 最小ACUをゼロより高く設定して常時応答性を保ち、最大ACUを予算とピークに合わせれば性能・接続・コスト境界を両立できる。

4 正解が要件を満たす理由

Bは低負荷時に設定した最小ACU付近まで縮小し、発売時には最大ACUまで自動的かつ細粒度に拡張するため、ピーク用インスタンスの常時費用を避ける。Aurora PostgreSQL互換と分散ストレージを維持し、予測時刻に依存するDBクラス変更や接続を伴う置換を避けられる。最小値で通常性能と接続余力、最大値で予算と最大性能を明示できる。

5・6 各選択肢の検討

不適切な理由と、正解になり得る条件を対にして確認します。

A 不正解

ピーク性能と安定性は満たすが、85%の低負荷時間にも最大DBクラスを払い続けるため、明示されたコスト削減条件に反する。

この選択肢が正解になる条件

負荷が常に高く予測可能で、最大性能の一貫性やプロビジョンドインスタンスの割引利用が自動縮小より重要なら有力である。

B 正解

ACUの最小・最大範囲内で変動負荷へ自動追従し、Aurora互換性と可用性を保ちながら低負荷時の過剰容量を減らす。

C 不正解

発売開始時刻が週ごとに異なるため固定スケジュールはピークを外し得る。DBクラス変更は細粒度な連続スケーリングではなく、フェイルオーバーや接続再確立を伴い得るため要件に合わない。

この選択肢が正解になる条件

負荷時刻が正確に予測でき、変更時間帯のフェイルオーバーや接続再確立を許容し、少数の既定DBクラスへ計画的に切り替える運用を受け入れられるなら候補になる。

D 不正解

Single-AZは可用性を低下させ、RDS DBインスタンスをEC2 Auto Scalingで増減したり同じEBSを共有マウントしたりする構成は成立しない。

この選択肢が正解になる条件

小規模な非本番DBでAZ障害を許容し、固定低負荷・最低費用が最優先ならSingle-AZ RDS自体は候補になる。

7 今後使える判断基準

Auroraの負荷が短時間・不規則・振幅大ならServerless v2の最小/最大ACU、定常高負荷ならプロビジョンドを比較する。可用性モードと容量課金は別軸で判断する。

8 関連サービスとの比較

Aurora Serverless v2はACU単位の自動増減、プロビジョンドAuroraは予測可能な定常負荷と割引に適する。RDS Multi-AZはスタンバイ可用性、read replicaは読み取り拡張であり、writer容量の自動縮小とは目的が異なる。

QUESTION 051

問題51

SECURE

複合難問

単一選択

NATなしLambdaのAWSサービス私設接続

医療APIのLambda関数は、インターネット経路を持たない3つのプライベートサブネットからAuroraへ接続する。処理中にS3から設定を読み、DynamoDBへ監査レコードを書き、Secrets ManagerからDB認証情報を取得する必要がある。規制上、NAT Gateway、Internet Gateway、パブリックIPは使用できない。現在VPC接続後に3サービスへの呼び出しがタイムアウトする。可用性のため各AZから利用でき、不要な時間・データ処理料金を抑えたい。DNS名をSDKで変更せず、実行ロールの最小権限も維持する。各サブネットのルートテーブルは個別で、Interface Endpointへ付けるセキュリティグループはVPC内のすべてではなく関数のSGだけを送信元にしたい。最も適切なネットワーク構成はどれか。

- A** 各AZにNAT Gatewayを作成し、LambdaサブネットのデフォルトルートにNATへ向ける。S3、DynamoDB、Secrets Managerの公開エンドポイントをHTTPSで呼ぶ。
- B** Lambda関数をVPCから外して3サービスへ接続し、Auroraにはパブリックアクセスを有効化してセキュリティグループでLambdaの送信元IPを許可する。
- C** S3とDynamoDBのGateway VPC Endpointを対象サブネットのルートテーブルに関連付ける。Secrets ManagerのInterface VPC Endpointを各AZに作りPrivate DNSを有効にし、エンドポイントSGでLambdaからの443を許可する。
- D** Secrets ManagerだけにGateway VPC Endpointを作り、S3とDynamoDBには単一AZのInterface EndpointをPrivate DNSなしで作る。SDKは従来の公開DNS名のまま使用する。

ここで答えを決める

正解・解説は次のページから

ANSWER 051

問題51の正解・解説

1 正解 C

- C** S3とDynamoDBのGateway VPC Endpointを対象サブネットのルートテーブルに関連付ける。Secrets ManagerのInterface VPC Endpointを各AZに作りPrivate DNSを有効にし、エンドポイントSGでLambdaからの443を許可する。

2 問題文の重要な条件

- VPC接続LambdaがS3、DynamoDB、Secrets ManagerへNATなしで到達する
- 3AZで利用可能にし、SDKのサービスDNS名を変更しない
- Gateway Endpointを使えるサービスでは不要なエンドポイント時間料金を避ける

3 正解へ至る判断過程

1. S3とDynamoDBはルートテーブル型のGateway Endpointを利用でき、NATなし・追加のエンドポイント時間料金なしでプライベート到達できる。
2. Secrets ManagerはPrivateLinkのInterface Endpointを使い、利用する各AZのサブネット、セキュリティグループ、Private DNSを正しく構成する。
3. ネットワーク到達性を作ってもAPI認可は別なので、Lambda実行ロールとエンドポイントポリシーを必要なリソースへ限定する。

4 正解が要件を満たす理由

Cは、S3とDynamoDBをGateway Endpointのプレフィックスリスト経路へ振り分け、Secrets Managerを各AZのプライベートENIへ名前解決する。いずれもインターネットやNATを必要とせず、通常のSDKエンドポイント名をPrivate DNSで維持できる。Gateway Endpointを使うサービスではInterface Endpointの時間料金を避け、Secrets ManagerだけをAZごとに配置して可用性を確保する。SG、endpoint policy、実行ロールで多層の最小権限も実装できる。

5・6 各選択肢の検討

不適切な理由と、正解になり得る条件を対にして確認します。

A 不正解

技術的には到達できるが、NATとインターネット経路の使用禁止に明確に反し、3AZ分の時間料金と処理料金も増える。

この選択肢が正解になる条件

外部インターネットAPIやVPC Endpoint非対応サービスへも広く接続する必要があり、NAT経路が規制上許可される場合は各AZ NATが適する。

B 不正解

VPC外LambdaからプライベートAuroraへ接続できず、Aurora公開化は攻撃面を広げる。Lambdaの固定送信元IPを前提とする許可も適切でない。

この選択肢が正解になる条件

関数がVPC内リソースへ一切接続せず、AWSの公開サービスエンドポイントだけを呼ぶなら、VPCへ接続しないLambdaは簡素で適切である。

C 正解

サービスごとに適切なGateway/Interface Endpointを選び、NATなし、3AZ、DNS互換、低コスト、最小権限を満たす。

D 不正解

Secrets ManagerにGateway Endpointはなく、Private DNSなしのInterface Endpointでは通常の公開DNS名が自動的に私設IPへ解決されない。単一AZも可用性条件に弱い。

この選択肢が正解になる条件

対象サービスがGateway Endpoint対応で、またはInterface Endpoint固有DNS名をアプリへ明示でき、単一AZ障害を許容する開発環境なら一部の考え方は成立する。

7 今後使える判断基準

NATなしでAWS APIを呼ぶときは、まずS3/DynamoDBのGateway Endpoint、その他対応サービスのInterface Endpointを判定する。InterfaceはAZ配置・SG・Private DNS、両者共通でendpoint policyとIAMを確認する。

8 関連サービスとの比較

Gateway EndpointはS3/DynamoDB向けのルートテーブル型で追加料金がなく、Interface EndpointはPrivateLink ENI型でAZごとの料金とSGがある。NATは幅広い外部宛先に使えるが公開経路と処理料金を伴う。

QUESTION 052

問題52

RESILIENT

本番相当

複数選択 (2つ)

フィルタ付きSNS/SQSファンアウト

ECサイトは注文確定イベントを、在庫、請求、分析、不正検知の4サービスへ配信する。各サービスは独立してスケールし、請求が2時間停止しても他の3サービスは処理を続け、復旧後に請求だけが滞留分を再処理できなければならない。イベントには国、支払方法、注文種別の属性があり、不正検知は特定支払方法、分析は特定国だけを受け取りたい。プロデューサーに4宛先を意識させず、各コンシューマーの一時障害と毒メッセージを隔離する。順序は不要で、少なくとも1回配信を前提に冪等処理する。各サービスの処理速度と保守時間は異なり、分析の滞留を解消するためにワーカーを増やしても請求サービスの可視性や再試行設定へ影響してはならない。選択すべき構成要素を2つ選択してください。

- A** プロデューサーは標準SNSトピックへ1回発行し、各SQSサブスクリプションにメッセージ属性または本文に基づくfilter policyを設定する。
- B** 各コンシューマー専用の標準SQSキューをSNSへ購読させ、キューごとに可視性タイムアウト、再試行回数、DLQ、滞留アラームを設定する。
- C** 4サービスが1つの標準SQSキューを共有し、競合コンシューマーとして受信する。不要なイベントを受け取ったサービスは削除せず可視性を戻す。
- D** SNSから4つのLambda関数を直接同期呼び出しし、失敗した関数がある間はプロデューサーも同じ注文を再送し続ける。永続キューとDLQは作らない。

ここで答えを決める

正解・解説は次のページから

ANSWER 052

問題52の正解・解説

1 正解 **A・B**

- A** プロデューサーは標準SNSトピックへ1回発行し、各SQSサブスクリプションにメッセージ属性または本文に基づくfilter policyを設定する。
- B** 各コンシューマー専用の標準SQSキューをSNSへ購読させ、キューごとに可視性タイムアウト、再試行回数、DLQ、滞留アラームを設定する。

2 問題文の重要な条件

- 1イベントを複数サービスへ複製し、プロデューサーを宛先から分離する
- 属性によって各サービスが必要なイベントだけを受け取る
- あるサービスの停止・毒メッセージ・滞留を他サービスから隔離する

3 正解へ至る判断過程

1. 複数購読者へ同じイベントを複製する役割はSNS、各購読者の永続バッファと独立再試行はSQSに分ける。
2. SNS subscription filter policyを使えば、プロデューサーやコンシューマーに不要イベントの破棄ロジックを増やさず選別できる。
3. キューをサービスごとに分けることで、滞留、可視性、DLQ、スケーリングの障害領域を独立させる。

4 正解が要件を満たす理由

Aはプロデューサーを単一SNSトピックに結合し、ファンアウトとフィルタリングを配信層へ移す。Bはサービスごとのキューに同じイベントの独立コピーを保持するため、請求停止中も他サービスが影響を受けず、請求だけが後で再開できる。各キューのDLQとCloudWatchアラームで毒メッセージと滞留も個別に管理でき、標準SNS/SQSの配信特性は冪等コンシューマーで補完される。

5・6 各選択肢の検討

不適切な理由と、正解になり得る条件を対にして確認します。

A 正解

一度の発行を複数購読へ複製し、subscription filter policyでサービス別の選択配信を実現する。

B 正解

コンシューマーごとの永続キュー、再試行、DLQ、監視により、停止と失敗を相互に隔離する。

C 不正解

1つのSQSメッセージは競合する受信者の一つへ渡るため、4サービスすべてへの複製にならない。不要メッセージの可視性戻しは再試行ループを作る。

この選択肢が正解になる条件

複数ワーカーが同じ種類のジョブを分担し、各メッセージをどれか1ワーカーだけが処理すればよいワークキューなら共有SQSが適する。

D 不正解

直接購読だけでは2時間停止分をコンシューマーごとのキューに保持できず、失敗処理をプロデューサーへ結合する。毒メッセージの隔離もない。

この選択肢が正解になる条件

処理が短く、失敗時のSNS再試行とDLQで十分で、コンシューマー独自の長い滞留や流量平準化が不要ならLambda直接購読も有効である。

7 今後使える判断基準

ファンアウトではSNSを複製・選別、SQSを購読者ごとの永続バッファ・再試行境界として使う。『全員へ届ける』ならキューを共有せず、1購読者1キューを基本にする。

8 関連サービスとの比較

SNSはpush型の多宛先配信とfilter policy、SQSはpull型の滞留・可視性・DLQを提供する。単一共有SQSは競合ワーカー、EventBridgeはイベントルーティングと多様なターゲット条件が中心である。

QUESTION 053

問題53

PERFORMANCE

本番より難しい

単一選択

既存SFTP利用者のマネージド移行

物流会社は取引先700社から毎日SFTPで配送ファイルを受信する。各社は既存の `sftp.partner.example.com`、ユーザー名、SSH公開鍵を使用しており、クライアント変更には半年以上かかる。現在の2台のLinux SFTPサーバーではパッチ、容量増設、HA切替が負担になっている。新環境では受信ファイルをS3へ直接保存し、特定のワークフローだけ共有POSIXファイル操作が必要なため別サーバーではEFSを使いたい。既存の企業ディレクトリまたはカスタム認証を連携し、CloudWatchへ監査ログを集約する。公開インターネットから到達するエンドポイントは必要だが、利用者ごとのホームディレクトリを分離し、保存先へのIAM権限を最小化する。SFTPプロトコルを維持し、サーバーOS運用をなくす最適なサービスはどれか。

- A** DataSyncエージェントを各取引先へ配布し、SFTP接続をDataSyncタスクへ置き換える。取引先はAWSコンソールからタスクを毎日開始する。
- B** Auto ScalingグループのEC2へOpenSSHを構成し、EFSへユーザー情報を保存する。Elastic IPを各インスタンスへ付け替え、OSパッチとfail2banを自社管理する。
- C** S3 presigned URLを毎日メールで取引先へ送り、全取引先にSFTPクライアントをHTTP PUT対応へ改修してもらう。EFSが必要なファイルは手動コピーする。
- D** AWS Transfer FamilyでSFTP対応サーバーを作成し、S3またはEFSをストレージを選ぶ。カスタムホスト名をDNSで割り当て、サービス管理、Directory Service、またはカスタムIdPで利用者を認証し、ログを有効化する。

ここで答えを決める

正解・解説は次のページから

ANSWER 053

問題53の正解・解説

1 正解 **D**

- D** AWS Transfer FamilyでSFTP対応サーバーを作成し、S3またはEFSをストレージを選ぶ。カスタムホスト名をDNSで割り当て、サービス管理、Directory Service、またはカスタムIdPで利用者を認証し、ログを有効化する。

2 問題文の重要な条件

- 700社の既存SFTPクライアント、ホスト名、SSH鍵／認証方式を極力維持する
- 保存先としてS3またはEFSを選び、監査ログを集約する
- SFTPサーバーOS、パッチ、HAの自己運用をなくす

3 正解へ至る判断過程

1. 要件は一括データ移行ではなく、外部利用者へ継続提供するマネージドSFTPエンドポイントである。
2. Transfer FamilyはSFTPプロトコルを終端し、バックエンドにS3またはEFS、認証に複数のIdP方式を選ぶ。
3. カスタムホスト名と既存資格情報の移行を設計すれば、取引先クライアントを改修せずOS運用をサービスへ移せる。

4 正解が要件を満たす理由

DはAWSが管理するSFTPサーバーを提供し、サーバーごとにS3またはEFSをドメインとして利用できる。既存DNS名をカスタムホスト名へ向け、サービス管理ユーザーの公開鍵、Directory Service、またはLambda/API Gatewayを使うカスタムIdPで認証を移せる。CloudWatchの構造化ログ等で接続・転送を監査でき、EC2のOS、SSHデーモン、パッチ、可用性を運用する必要がない。

5・6 各選択肢の検討

不適切な理由と、正解になり得る条件を対にして確認します。

A 不正解

DataSyncは管理されたストレージ間転送用で、700社へ公開する対話型SFTPサーバーや既存SFTP資格情報の終端サービスではない。

この選択肢が正解になる条件

自社管理のNFS/SMB/オブジェクトストレージ間で大量データを初回・増分転送し、双方の管理権限を持つ場合はDataSyncが適する。

B 不正解

既存SFTP互換性は高いが、OSパッチ、SSH設定、スケーリング、状態共有、IP切替を引き続き自社管理し、中心要件に反する。

この選択肢が正解になる条件

Transfer Familyで使えない独自SSHサブシステムやOS拡張が必須で、完全なホスト制御のため運用負荷を受け入れる場合はEC2が必要になる。

C 不正解

S3への直接アップロード自体は簡素だが、全取引先のプロトコルとクライアント変更が必要で、半年以上変更できない制約を満たさない。

この選択肢が正解になる条件

新規のWeb/モバイルクライアントだけが対象でHTTPアップロードへ変更でき、短時間のオブジェクト単位権限を発行したい場合はpresigned URLが適する。

D 正解

既存SFTP接続を保ちながらS3/EFS、複数認証方式、カスタムDNS、監査ログをマネージドに統合し、OS運用を除去する。

7 今後使える判断基準

外部利用者が既存SFTP/FTPS/FTPを維持しつつAWSストレージへ移るならTransfer Familyを検討する。DataSyncは管理者間の転送、presigned URLはプロトコル変更可能な直接オブジェクトアップロードである。

8 関連サービスとの比較

Transfer Familyはマネージドなファイル転送プロトコル終端、DataSyncはストレージ間コピー、EC2 OpenSSHは最大の自由度と最大の運用、presigned URLはS3 APIへの一時直接アクセスを提供する。

QUESTION 054

問題54

COST

複合難問

単一選択

共有ストレージ要件を外した大容量素材保管

広告会社は220TBの完成済み画像・動画を20台のLinuxレンダリングワーカーで参照する。オブジェクトの98%は作成後に変更されず、キーで取得するようアプリを改修できる。月間参照量は全体の8%で、90日後はさらに低下する。残り2%の一時作業領域だけはワーカー間のPOSIX共有とファイルロックが必要で、処理完了後7日で削除できる。Windowsクライアント、SMB、NTFS ACLは不要である。現在は全データを高性能共有NASに置くため費用が高い。S3へ移す完成素材には同時上書きやディレクトリrenameはなく、取り出し遅延を許容できる古い世代と即時取得が必要な新しい世代を区別できる。耐久性を保ち、アクセス料金と変更工数を考慮しながら総ストレージ費を最も下げる構成はどれか。

- A** 完成済み220TBの大部分をS3へオブジェクトとして保存し、アクセス低下に応じたLifecycleまたは適切なストレージクラスを使う。短命な2%のPOSIX作業領域だけをEFSに置く。
- B** 完成済みデータと作業データをすべてEFS Standardへ置き、各AZのマウントターゲットからNFSアクセスする。Lifecycleは使用せず常に同じ性能クラスに保つ。
- C** 全データをFSx for Windows File Server Multi-AZへ移し、LinuxワーカーにSMBクライアントを追加する。NTFS ACLとWindowsファイルロックで管理する。
- D** 各ワーカーに同じEBSボリュームをMulti-Attachしてext4でマウントし、220TBを複製せず共有する。スナップショットを唯一の耐久コピーとする。

ここで答えを決める

正解・解説は次のページから

ANSWER 054

問題54の正解・解説

1 正解 **A**

- A** 完成済み220TBの大部分をS3へオブジェクトとして保存し、アクセス低下に応じたLifecycleまたは適切なストレージクラスを使う。短命な2%のPOSIX作業領域だけをEFSに置く。

2 問題文の重要な条件

- 大部分は不変でキー取得でき、共有ファイルシステムのセマンティクスを必要としない
- POSIX共有とロックが必要なのは短命な2%だけである
- 220TBの耐久性を保ちつつ、アクセス低下に応じて保存費を下げる

3 正解へ至る判断過程

1. すべてを同じ共有ファイルサービスへ置く前提を外し、アクセスセマンティクスが必要なデータだけにファイルシステム料金を支払う。
2. 不変・キー取得・大容量の完成素材はS3オブジェクトとストレージクラス／Lifecycleに適合する。
3. 小さな短期作業領域だけをEFSへ残すことで、POSIX共有・ロックを保ちながら高価な常時ファイルストレージ容量を最小化する。

4 正解が要件を満たす理由

Aは大容量の不変データを高耐久なS3へ移し、実際のアクセス頻度に応じてStandard-IA、Intelligent-Tiering、アーカイブ等を選択できる。アプリはキー取得へ変更可能なのでPOSIXセマンティクスを失っても支障がない。一方、処理中に本当に共有ロックが必要な2%だけをEFSへ置き、7日で削除すればファイルシステム容量と期間を大きく減らせる。要件に応じた二層化が総コストを最も下げる。

5・6 各選択肢の検討

不適切な理由と、正解になり得る条件を対にして確認します。

A 正解

オブジェクトとファイルの要件を分離し、大容量部分にS3の低い保存単価と階層化、小部分にEFSのPOSIX共有を適用する。

B 不正解

EFSは必要なPOSIX共有を満たすが、不変でファイルセマンティクス不要な220TBまでStandardに置くと過剰な機能と保存費を支払う。

この選択肢が正解になる条件

全データを多数Linuxホストが低遅延NFSで更新し、ディレクトリ操作、POSIX権限、ファイルロックを常時必要とするならEFSが適する。

C 不正解

FSx for WindowsはSMB/Windows互換要件に強いが、本件には不要で、Linux中心の不変オブジェクトにMulti-AZ Windowsファイルサーバー費用を加える。

この選択肢が正解になる条件

利用者がWindowsでSMB、NTFS ACL、Active Directory統合を必須とし、既存アプリを変更できない場合はFSx for Windowsが適する。

D 不正解

EBSは基本的にブロックストレージで、通常のext4を複数インスタンスから安全な共有ファイルシステムとして同時マウントできない。AZや耐久設計も要件に合わない。

この選択肢が正解になる条件

単一EC2に低遅延ブロックストレージが必要な場合、または対応io1/io2とクラスタ対応ファイルシステムを限定的に運用する場合はEBSを選べる。

7 今後使える判断基準

ストレージはデータ全体ではなくデータ集合ごとのアクセスセマンティクスで選ぶ。不変・API/キー取得=S3、Linux POSIX共有=EFS、Windows SMB/NTFS=FSx for Windowsとし、共有要件のない容量をファイルシステムへ置かない。

8 関連サービスとの比較

S3は大規模オブジェクトと階層化、EFSはRegionalなNFS/POSIX共有、FSx for WindowsはSMB/NTFS/AD、EBSはEC2向けブロックストレージであり、同じ『保存』でもインターフェイスと課金単位が異なる。

QUESTION 055

問題55

SECURE

本番相当

単一選択

VPC内限定のAPI Gateway REST API

銀行はAPI Gatewayで勘定照会REST APIを公開する。呼出元は同一リージョンの3つのVPCにあるEC2とLambdaだけで、インターネット、NAT、パブリックAPIエンドポイント経由のアクセスは禁止する。各VPCから通常のexecute-api DNS名を利用し、TLSで接続したい。特定のVPC Endpoint以外からはAPIを呼べず、許可したIAMロールだけがメソッドを実行できるよう二重に制限する。バックエンドはVPC内のサービスだが、APIの公開面をALBそのものにはしたくない。3つのVPCは同一Organizations内だが運用主体が異なるため、VPC CIDRだけでなく呼出経路とIAM principalの両方を監査証跡に残したい。最小の公開範囲で要件を満たす構成はどれか。

- A** API GatewayのREGIONAL REST APIを作り、AWS WAFで3つのVPC CIDRを許可する。各VPCはNAT Gateway経由でパブリックexecute-apiエンドポイントへ接続する。
- B** API GatewayのPRIVATE REST APIを作り、各VPCにexecute-api用Interface VPC Endpointを配置してPrivate DNSを有効にする。API resource policyをvpce IDで限定し、IAM認証とendpoint policyも必要最小限にする。
- C** インターネット向けALBへAPIを直接公開し、セキュリティグループの送信元に3つのVPC CIDRを指定する。API Gatewayは使用しない。
- D** API Gatewayからバックエンドへ接続するVPC Linkだけを作成する。API自体はEDGE最適化のままにし、CloudFrontの地理的制限で社内VPCだけを識別する。

ここで答えを決める

正解・解説は次のページから

ANSWER 055

問題55の正解・解説

1 正解 **B**

- B** API GatewayのPRIVATE REST APIを作り、各VPCにexecute-api用Interface VPC Endpointを配置してPrivate DNSを有効にする。API resource policyをvpce IDで限定し、IAM認証とendpoint policyも必要最小限にする。

2 問題文の重要な条件

- APIの呼出経路をVPC内のPrivateLinkに限定し、公開エンドポイントを使わない
- 通常のexecute-api名を使うためPrivate DNSが必要である
- vpce単位のresource policyとIAM主体単位の認証を組み合わせる

3 正解へ至る判断過程

1. API Gatewayのフロントドアを非公開にする要件なので、バックエンド接続用VPC LinkではなくPRIVATE APIエンドポイントタイプを選ぶ。
2. クライアントVPCにはexecute-apiのInterface Endpointを作り、Private DNSでSDKや既存URLをプライベートIPへ解決する。
3. resource policyで許可VPC Endpointを限定し、endpoint policyとIAM認証で通過主体とAPI操作をさらに絞る。

4 正解が要件を満たす理由

BはPrivate REST APIをInterface VPC Endpoint経由でのみ到達可能にし、トラフィックをAWSネットワーク内に保つ。Private DNSにより既存のexecute-apiホスト名を変えずにエンドポイントENIへ接続できる。APIのresource policyでaws:SourceVpce等を使って許可経路を限定し、VPC endpoint policyとIAM認証で利用者・APIを絞るため、ネットワーク境界とアイデンティティ境界を両方満たす。

5・6 各選択肢の検討

不適切な理由と、正解になり得る条件を対にして確認します。

A 不正解

REGIONAL APIは公開エンドポイントで、NATとインターネット経路禁止に反する。WAFのCIDR制限もVPC Endpointそのものの強制にはならない。

この選択肢が正解になる条件

外部顧客も利用する公開APIで、WAF、認証、レート制限を使ってインターネットから保護する場合はREGIONAL APIが適する。

B 正解

Private API、Interface Endpoint、Private DNS、resource/IAM/endpoint policyにより、非公開経路と多層最小権限を実現する。

C 不正解

インターネット向けALBは公開面を作り、API Gatewayの認可、使用量制御、メソッド機能も失う。CIDRだけでは指定vpce経由を保証しない。

この選択肢が正解になる条件

API Gateway機能が不要で、VPC内クライアント向けのinternal ALBとしてHTTPサービスを直接提供するならALBだけの構成が簡素である。

D 不正解

VPC LinkはAPI Gatewayからプライベートバックエンドへの接続であり、クライアントからAPI Gatewayへの公開経路を非公開化しない。地理制限もVPC識別ではない。

この選択肢が正解になる条件

公開API GatewayからVPC内のNLB/ALB等へプライベート統合することだけが要件ならVPC Linkを使用する。

7 今後使える判断基準

API Gateway問題では『クライアント→API』と『API→バックエンド』を分ける。前者を私設化するのはPRIVATE API+execute-api Interface Endpoint、後者を私設化するのはVPC Linkである。

8 関連サービスとの比較

PRIVATE REST APIはVPC Endpoint経由の社内API、REGIONAL/EDGE APIは公開クライアント向け、VPC LinkはAPI GatewayからVPCバックエンドへの統合、internal ALBはAPI管理機能不要のVPC内HTTP負荷分散向けである。

QUESTION 056

問題56

RESILIENT

本番より難しい

複数選択 (2つ)

JMS互換を維持するマネージドブローカー移行

航空会社はオンプレミスのActiveMQで、JMS 1.1アプリケーション40本を接続している。JMS トランザクション、selectors、durable subscriptions、既存のOpenWire接続を多用しており、6か月以内にアプリケーションを書き換える余裕がない。AWS移行後はブローカーのOSとパッチをAWSへ委ね、1つのAZまたはブローカーホスト障害から自動復旧したい。メッセージは保管時・転送時に暗号化し、接続断時はクライアントが別エンドポイントへ再接続する。クラウドネイティブサービスへの再設計は次期計画とする。現行クライアントはActiveMQの failover URIを設定でき、DNS切替だけに依存せず複数の接続先を試行する機能を持っている。選択すべき対応を2つ選択してください。

- A** Amazon MQ for ActiveMQをactive/standbyのMulti-AZブローカーとして作成し、既存JMS/OpenWire互換性とマネージド運用を利用する。
- B** すべての宛先を標準SQSとSNSに置き換え、JMSトランザクションとselectorsを同じ設定ファイルのまま自動変換する。アプリコードは変更しない。
- C** クライアントを複数ブローカーエンドポイントへ接続できるfailover transportとTLSに設定し、宛先・ユーザー・設定を検証して段階的にメッセージを移行する。
- D** Amazon MSKクラスターを作成し、既存JMSクライアントのOpenWire URLだけをKafka bootstrap serverへ変更する。durable subscriptionはKafkaが透過的に再現する。

ここで答えを決める

正解・解説は次のページから

ANSWER 056

問題56の正解・解説

1 正解 A・C

- A** Amazon MQ for ActiveMQをactive/standbyのMulti-AZブローカーとして作成し、既存JMS/OpenWire互換性とマネージド運用を利用する。
- C** クライアントを複数ブローカーエンドポイントへ接続できるfailover transportとTLSに設定し、宛先・ユーザー・設定を検証して段階的にメッセージを移行する。

2 問題文の重要な条件

- 既存ActiveMQのJMS/OpenWire機能を大幅なコード変更なしで維持する
- ブローカーOS運用を委ね、AZまたはホスト障害に自動フェイルオーバーする
- クライアント側も冗長エンドポイント、TLS、段階移行へ対応する

3 正解に至る判断過程

1. 短期の最優先条件はクラウドネイティブ化ではなく、ActiveMQ/JMSのセマンティクス互換性である。
2. Amazon MQ for ActiveMQは既存プロトコルとAPIを保ちながら管理型ブローカーへ移せ、active/standby構成で冗長化できる。
3. サービス側を冗長化してもクライアントが単一エンドポイント固定では復旧できないため、failover transportと移行検証を組み合わせる。

4 正解が要件を満たす理由

AはActiveMQ互換のAmazon MQを選ぶため、JMSの機能やOpenWire接続を保ちながら、ブローカーソフトウェアと基盤運用をAWSへ移せる。active/standby Multi-AZは別AZのスタンバイへ障害時に切り替える。Cはクライアントに冗長エンドポイントを認識させ、TLSで保護し、切替時に自動再接続できるようにする。設定・宛先・未処理メッセージを段階検証することで、互換性だけでは解決しない移行リスクも抑える。

5・6 各選択肢の検討

不適切な理由と、正解になり得る条件を対にして確認します。

A 正解

ActiveMQ/JMS/OpenWire互換を維持しつつ、管理されたactive/standby Multi-AZブローカーで運用と障害対応を改善する。

B 不正解

SQS/SNSは有力なマネージドメッセージングだが、JMSトランザクション、selectors、durable subscriptionsを無変更で透過置換するサービスではない。

この選択肢が正解になる条件

アプリを非同期AWS APIへ再設計でき、JMS固有機能を不要にし、サーバーレスなキューとファンアウトを優先する次期移行なら適する。

C 正解

Multi-AZブローカーの冗長エンドポイントをクライアントが実際に利用し、暗号化と段階移行を成立させる。

D 不正解

KafkaとActiveMQはプロトコルと配信モデルが異なり、OpenWire URLの変更だけでJMS機能を透過的に移行できない。大規模なアプリ改修が必要になる。

この選択肢が正解になる条件

イベントストリーミングへ再設計でき、高スループットのログ、パーティション、コンシューマーグループを活用する要件ならMSKが適する。

7 今後使える判断基準

メッセージング移行は『API/プロトコル互換性を維持する期限』があるかで分ける。JMS互換の短期リホストはAmazon MQ、再設計可能ならSQS/SNSやMSKを配信モデルに応じて選ぶ。HAはクライアント再接続設定まで確認する。

8 関連サービスとの比較

Amazon MQはActiveMQ/RabbitMQ互換の管理ブローカー、SQS/SNSはサーバーレスなキュー／ファンアウト、MSKはKafkaストリーミングである。Amazon MQのactive/standbyだけでなく failover transportがエンドツーエンド復旧に必要となる。

QUESTION 057

問題57

PERFORMANCE

複合難問

単一選択

S3データを処理するLustre HPC共有

創薬企業はS3にある500TBのゲノムデータを、毎月3日間だけ200台のEC2から並列解析する。各ノードは同じデータセットをPOSIXパスで参照し、合計数百GB/秒規模のスループットと高い並列I/Oが必要である。入力の本拠はS3にあり、処理中間データは再生成可能だが、最終結果はジョブ完了時にS3へ戻して長期保管する。月の残り期間に高性能ファイルシステム料金を払い続けたくなく、各ノードへ500TBを個別コピーする時間も避けたい。解析ジョブは一斉に同じ参照ファイルを読む段階と、大量の一時ファイルを並列生成する段階があり、メタデータ操作も集中的に発生する。S3とのデータ連携と共有名前空間を保ちながら、HPC性能とコストを最も適切に両立する構成はどれか。

- A** EFS RegionalをGeneral Purposeモードで作成し、S3全データを毎月cpでコピーする。200台が同じNFS共有を使用し、処理後も翌月まで500TBをEFS Standardに保持する。
- B** 各EC2に500TB相当のEBS gp3ボリュームを割り当て、S3から全入力を個別ダウンロードする。最終結果は各ノードのスナップショットとして保存する。
- C** ジョブ期間にFSx for Lustreを作成してS3バケットとData Repository Associationで関連付ける。並列POSIXアクセスで処理し、結果をS3へエクスポートしてからファイルシステムを削除する。
- D** オンプレミスにS3 File Gatewayを作成し、200台のEC2からDirect Connect越しにNFS共有をマウントする。ゲートウェイの単一ローカルキャッシュで全並列I/Oを処理する。

ここで答えを決める

正解・解説は次のページから

ANSWER 057

問題57の正解・解説

1 正解 C

- C** ジョブ期間にFSx for Lustreを作成してS3バケットとData Repository Associationで関連付ける。並列POSIXアクセスで処理し、結果をS3へエクスポートしてからファイルシステムを削除する。

2 問題文の重要な条件

- 200台から同一のPOSIX名前空間へ非常に高い並列I/Oが必要である
- S3が入力の正本かつ結果の長期保管先で、中間データは再生成できる
- 月3日だけファイルシステムを利用し、ノードごとの全量コピーと常時費用を避ける

3 正解へ至る判断過程

1. 通常の共有NFSではなく、HPCノードへデータをストライプして高い集約スループットを提供する並列ファイルシステムが必要である。
2. FSx for LustreはS3データリポジトリと関連付け、S3オブジェクトをファイル名前空間から処理し、結果を戻せる。
3. 正本をS3へ残し、必要な期間だけLustreを作成してエクスポート後に削除すれば、性能用ストレージの常時費用を避けられる。

4 正解が要件を満たす理由

CはLustreの分散メタデータとストライピングされたストレージにより、多数EC2から同一POSIX名前空間へ高スループットでアクセスできる。S3とのData Repository Associationにより入力をファイルとして見せ、必要時に取り込み、生成結果を自動エクスポートまたはデータリポジトリタスクでS3へ戻せる。中間データが再生成可能なので一時的な処理用途に合わせたデプロイを選び、結果確認後に削除することで月3日分に費用を集中できる。

5・6 各選択肢の検討

不適切な理由と、正解になり得る条件を対にして確認します。

A 不正解

EFSは汎用NFS共有として有効だが、数百GB/秒級のHPC並列I/OとLustre最適化を前提とせず、500TBを常時Standardに保持する費用も条件に合わない。

この選択肢が正解になる条件

多数Linuxアプリが日常的に汎用NFS、POSIX権限、伸縮容量を必要とし、HPC級の一時的集約スループットが不要ならEFSが適する。

B 不正解

ノードごとの500TB複製は時間・容量・転送を200倍にし、共有名前空間も失う。EBSスナップショットは処理結果の共有オブジェクト保管にも不向きである。

この選択肢が正解になる条件

各ノードが独立した小さなデータ分割だけを処理し、共有ファイルアクセスがなく、低遅延ブロックI/Oが必要ならノード別EBSは有効である。

C 正解

HPC向け並列POSIX性能、S3との入出力連携、期間限定作成によるコスト最適化を一つの構成で満たす。

D 不正解

File GatewayはハイブリッドなS3ファイルアクセス用で、単一オンプレミスキャッシュと回線が200台のリージョン内HPC I/Oのボトルネックになる。

この選択肢が正解になる条件

オンプレミスの少数アプリからS3をNFS/SMB共有として利用し、ホットデータのローカルキャッシュが必要ならFile Gatewayが適する。

7 今後使える判断基準

S3を正本に、多数ノードが同じPOSIXパスへHPC級並列I/Oを行うならFSx for Lustreを検討する。再生成可能な一時処理はLustreを期間限定にし、永続結果だけをS3へ戻す。

8 関連サービスとの比較

FSx for LustreはHPC並列ファイルとS3連携、EFSは汎用NFS共有、EBSは単一/限定クラスター向けブロック、S3 File Gatewayはオンプレミスのハイブリッドファイルアクセス向けである。

QUESTION 058

問題58

COST

本番相当

単一選択

定常RDSと可変コンピュートの購入分離

SaaS企業は、RDS for MySQL Multi-AZのdb.r6i系を2年間ほぼ一定サイズで24時間利用する。一方、アプリ層はEC2、Fargate、Lambdaを混在させ、今後1年でx86からGraviton、EC2からFargate、別リージョンへの一部移行を予定している。アプリ層の最低利用額は安定しているが、実行サービス、インスタンスファミリー、OSは変わる。繁忙期の超過分は3か月だけ発生する。割引契約の前に過去6か月の時間別使用量を確認し、組織の連結請求内で確実に消費されるベースラインだけを算出している。データベースの割引と、変更が多いコンピュートの柔軟性を両立し、未使用コミットを避ける購入戦略はどれか。

- A** RDSとアプリ層の両方に同じEC2 Standard Reserved Instanceを購入する。FargateとLambdaにもEC2 RIの割引が自動適用される前提で全ピークをコミットする。
- B** 全費用を1つのCompute Savings Planでコミットし、RDS DBインスタンス時間にも同じCompute Savings Plansの割引が適用されるものとして2年分を購入する。
- C** 構成変更の可能性があるため、RDSの定常利用を含む全リソースをOn-Demandのままにする。繁忙期だけゾーンReserved Instanceを購入して後で返却する。
- D** 安定したRDSベースラインには対応するReserved DB Instanceを購入する。アプリ層の安定した最低利用額だけにCompute Savings Plansを適用し、短期ピークはOn-Demand、耐障害ジョブはSpotで補う。

ここで答えを決める

正解・解説は次のページから

ANSWER 058

問題58の正解・解説

1 正解 **D**

- D** 安定したRDSベースラインには対応するReserved DB Instanceを購入する。アプリ層の安定した最低利用額だけにCompute Savings Plansを適用し、短期ピークはOn-Demand、耐障害ジョブはSpotで補う。

2 問題文の重要な条件

- RDS DBインスタンスは同じエンジン・クラス系で長期に安定して稼働する
- アプリ層はEC2/Fargate/Lambda、アーキテクチャ、リージョンが変わる
- コミットは安定ベースラインだけに限定し、3か月のピークを過剰購入しない

3 正解へ至る判断過程

1. RDSの割引商品とEC2/Fargate/Lambdaのコンピュータ割引を同一商品とみなさず、サービスごとに安定性を評価する。
2. 定常RDSにはReserved DB Instanceの割引が適し、Multi-AZを含む対応属性と期間を実利用へ合わせる。
3. 変更が多いアプリ層には広い柔軟性を持つCompute Savings Plansを最低利用額だけコミットし、変動分を従量／中断可能料金へ残す。

4 正解が要件を満たす理由

Dは、予測可能なRDS利用と変更が多いコンピュータ利用を別々の購入モデルで最適化する。Reserved DB Instanceは対象RDSの定常時間へ割引を適用し、Compute Savings Plansはコミット額の範囲でEC2、Fargate、Lambdaやファミリー・リージョン変更へ柔軟に追従する。ピークまでコミットせずOn-Demandで補い、中断可能な処理だけSpotへ回すため、割引率だけでなく未使用コミットのリスクも抑えられる。

5・6 各選択肢の検討

不適切な理由と、正解になり得る条件を対にして確認します。

A 不正解

EC2 RIはRDS Reserved DB Instanceとは別商品で、FargateやLambdaへ自動適用されない。ピーク全量のStandard RIは構成変更と未使用リスクも大きい。

この選択肢が正解になる条件

対象が同一リージョン・属性の予測可能なEC2だけで、長期間インスタンス構成を固定するならEC2 RI自体は割引候補になる。

B 不正解

Compute Savings Plansは柔軟なコンピュート利用向けだが、RDS Reserved DB Instanceの代わりに同じ割引をRDSへ適用する商品ではない。全費用コミットも過剰である。

この選択肢が正解になる条件

費用の対象がEC2、Fargate、LambdaなどCompute Savings Plans対応コンピュートだけで、安定した最低利用額に限定するなら適切である。

C 不正解

可変部分をOn-Demandにする考えはよいが、2年安定するRDSの割引機会を失う。Reserved Instanceは短期の繁忙期だけ購入して自由返却する商品でもない。

この選択肢が正解になる条件

全利用が数週間だけで将来量が全く読めず、コミットを使い切れないリスクが割引を上回る場合は全On-Demandが合理的である。

D 正解

RDS固有の定常割引と、サービス横断で変化するアプリの柔軟な割引を分離し、ピークへの過剰コミットを避ける。

7 今後使える判断基準

購入割引は『どのサービスに適用されるか』『どこまで利用が安定するか』で分ける。RDSはDB向け予約、EC2/Fargate/Lambdaの可変構成はCompute Savings Plansをベースラインだけに適用し、ピークは従量で残す。

8 関連サービスとの比較

Reserved DB InstancesはRDSエンジン・リージョン等に紐づくDB割引、Compute Savings PlansはEC2/Fargate/Lambdaをまたぐ柔軟なコンピュータ割引、On-Demandは変動分、Spotは中断可能処理向けである。

QUESTION 059

問題59

SECURE

本番より難しい

単一選択

CloudFrontとALBで異なるACMリージョン

欧州企業はwww.example.comをCloudFrontで配信し、オリジンとしてeu-west-1のALBを使用する。利用者からCloudFrontまでと、CloudFrontからALBまでの両区間をHTTPSにし、どちらもexample.com系の企業ドメインで証明書検証する。証明書はACMで自動更新し、秘密鍵をEC2へ配置しない。担当者はeu-west-1でワイルドカード証明書を1枚発行したが、CloudFrontのViewer certificate一覧に表示されない。ALBはeu-west-1に固定される。CloudFrontのalternate domain nameとALBオリジン用DNS名は異なるが、両方とも発行する証明書のSANへ含めてドメイン所有を検証できる。最小の変更で正しく証明書を配置する方法はどれか。

- A** CloudFrontのビューワーHTTPS用証明書をus-east-1で発行またはインポートし、ALBのHTTPSリスナー用証明書はeu-west-1で用意する。必要な各ドメイン名を両証明書のSANでカバーする。
- B** eu-west-1の既存ACM証明書1枚をCloudFrontとALBの両方へ関連付ける。CloudFrontはオリジンリージョンから証明書を自動複製するまで待つ。
- C** us-east-1で証明書を1枚だけ発行し、CloudFrontとeu-west-1のALBへ同じARNを設定する。ALBはグローバルACM ARNをどのリージョンからでも参照できる。
- D** 各EC2へ自己署名証明書を配置し、ALBをTCPパススルーに変更する。CloudFrontではデフォルトcloudfront.net証明書だけを使い、www.example.comの検証を省略する。

ここで答えを決める

正解・解説は次のページから

ANSWER 059

問題59の正解・解説

1 正解 **A**

- A** CloudFrontのビューワーHTTPS用証明書をus-east-1で発行またはインポートし、ALBのHTTPSリスナー用証明書はeu-west-1で用意する。必要な各ドメイン名を両証明書のSANでカバーする。

2 問題文の重要な条件

- CloudFrontのカスタムドメインに対するビューワー証明書が必要である
- eu-west-1のALBにもCloudFrontオリジン接続用HTTPS証明書が必要である
- ACM証明書はリージョナルリソースで、CloudFrontだけ特定リージョン要件がある

3 正解へ至る判断過程

1. CloudFrontはグローバルサービスだが、カスタムビューワー証明書はACMのus-east-1に存在する必要がある。
2. ALBはリージョナルリソースで、リスナーへ関連付けるACM証明書はALBと同じeu-west-1に必要である。
3. したがって同じドメインを含む証明書を二つのリージョンで管理し、各TLS区間へ別ARNを割り当てる。

4 正解が要件を満たす理由

AはCloudFront固有のus-east-1要件とALBの同一リージョン要件を同時に満たす。CloudFrontはus-east-1のACM証明書をエッジへ配布し、eu-west-1のALBは同リージョンの証明書でオリジンTLSを終端する。ACMのDNS検証等を使えば両証明書を管理型で更新でき、EC2へ秘密鍵を配布する必要もない。

5・6 各選択肢の検討

不適切な理由と、正解になり得る条件を対にして確認します。

A 正解

ビューワー側をus-east-1、リージョナルALB側をeu-west-1に置き、二つのTLS終端それぞれの証明書要件を満たす。

B 不正解

CloudFrontはeu-west-1のACM証明書をビューワー証明書として選択できず、オリジンリージョンから自動複製されることもない。

この選択肢が正解になる条件

CloudFrontを使わず、eu-west-1のALBだけでカスタムドメインHTTPSを終端する場合は既存の同リージョン証明書1枚でよい。

C 不正解

us-east-1証明書はCloudFrontには使えるが、eu-west-1のALBは別リージョンのACM証明書ARNをリスナーへ関連付けられない。

この選択肢が正解になる条件

オリジンもus-east-1のALBであれば同一リージョン内の証明書を利用できるが、CloudFront用とALB用の関連付け・名前はなお個別確認が必要である。

D 不正解

www.example.comをCloudFrontでHTTPS提供するにはその名前をカバーする信頼済み証明書が必要で、デフォルト証明書では代替できない。ALBもTCPパススルーを行わない。

この選択肢が正解になる条件

利用者がcloudfront.netのデフォルトドメインだけを使用しカスタムドメインを不要とする検証環境なら、CloudFrontデフォルト証明書を利用できる。

7 今後使える判断基準

ACM証明書はTLS終端リソースのリージョンで考える。CloudFrontのビューワー証明書はus-east-1、ALB/API Gateway Regionalなどのリージョナル終端はそのリソースと同じリージョンに置く。

8 関連サービスとの比較

CloudFrontはus-east-1のACM証明書をグローバル配布し、ALBは同一リージョンのACM証明書を使う。CloudFront→ALBとviewer→CloudFrontは別TLS区間なので、名前、発行リージョン、更新をそれぞれ確認する。

QUESTION 060

問題60

RESILIENT

複合難問

複数選択 (2つ)

APIの流量制御と非同期バッファ

画像変換サービスはAPI Gateway REST APIでジョブを受け付け、バックエンドDBは毎秒300件を超える更新で不安定になる。通常は毎秒80件だが、顧客の一括投入で10秒間に50,000件届く。クライアントは即時変換結果を必要とせず、ジョブを耐久的に受理できればHTTP 202と追跡IDを受け取れる。契約プランごとに目標要求レートと月間量をベストエフォートで抑制し、暴走クライアントには429と再試行指針を返したい。usage planを厳密な課金上限や認可の代替にはせず、超過分を絶対に受け付けない請求統制は別システムで行う。一方、入口で許可されたバーストは失わず後で処理し、ワーカー障害や毒メッセージを再試行・隔離する。受理と実処理は分離でき、同じ追跡IDの再送は二重ジョブを作らないよう冪等化する。DB保護と利用者公平性を両立するために選ぶべき対策を2つ選択してください。

- A** API Gateway REST APIでステージ／メソッドのthrottlingと、顧客APIキー別usage planのレート・バースト・クォータをベストエフォートの利用抑制として設定し、クライアントに429の指数バックオフを実装させる。
- B** API Gatewayから同期Lambdaを呼び、予約済み同時実行数を無制限相当に増やして全ジョブを直ちにDBへ書く。DBエラー時だけクライアントへ500を返す。
- C** API Gatewayの統合タイムアウトを最大化し、各HTTP接続をDB更新が完了するまで開いたままにする。キューを使わずクライアント接続をバッファとして扱う。
- D** 受理処理でジョブを標準SQSへ保存して追跡IDを返し、ワーカーの同時実行をDBの300件/秒以下に制御する。Visibility Timeout、再試行、DLQ、滞留メトリクスを設定する。

ここで答えを決める

正解・解説は次のページから

ANSWER 060

問題60の正解・解説

1 正解 **A・D**

- A** API Gateway REST APIでステージ／メソッドのthrottlingと、顧客APIキー別usage planのレート・バースト・クォータをベストエフォートの利用抑制として設定し、クライアントに429の指数バックオフを実装させる。
- D** 受理処理でジョブを標準SQSへ保存して追跡IDを返し、ワーカーの同時実行をDBの300件/秒以下に制御する。Visibility Timeout、再試行、DLQ、滞留メトリクスを設定する。

2 問題文の重要な条件

- REST APIで顧客別の目標レート・バースト・月間量をベストエフォートで入口抑制する
- 許可された50,000件バーストは耐久的に保持し、非同期処理する
- DB処理率を300件/秒以下に平準化し、失敗を再試行・隔離する

3 正解へ至る判断過程

1. 悪意・契約超過の要求と、正当だが下流能力を超えるバーストを同じ仕組みで扱わず、入口制限と内部バッファに分ける。
2. API Gateway REST APIのthrottling/usage planを公平性と過負荷抑制に使う。ただしクォータはベストエフォートであり、厳密な課金上限とは分ける。
3. 受理済みジョブはSQSへデカップリングし、コンシューマー同時実行を下流限界に合わせ、DLQと監視で障害を隔離する。

4 正解が要件を満たす理由

AはREST API全体の保護と顧客別の目標レート・クォータを入口でベストエフォート適用し、過剰要求を早期に抑える。これは厳密な請求上限や認可ではないという前提にも一致する。Dは受理可能なジョブをSQSへ耐久保存して短時間で202を返し、ワーカー側の同時実行を調整してDBへの流入を平準化する。Visibility Timeoutと再試行は一時障害、DLQは反復失敗を扱うため、バースト時にも損失と下流崩壊を避けられる。

5・6 各選択肢の検討

不適切な理由と、正解になり得る条件を対にして確認します。

A 正解

REST APIの利用者単位で目標レート、バースト、クォータをベストエフォート適用し、暴走要求と公平性を管理する。

B 不正解

Lambdaを大量並列化するとDBの300件/秒制約を直接超え、同期失敗をクライアント再送へ波及させる。耐久バッファもない。

この選択肢が正解になる条件

下流も自動的に同じ規模へ拡張でき、各処理が短く即時応答必須で、バッファや順序待ちが不要なら同期Lambdaの並列性を活用できる。

C 不正解

HTTP接続は耐久キューではなく、タイムアウト・再送・接続数を増やす。DB流量を安全に平準化せず、202非同期要件の利点も失う。

この選択肢が正解になる条件

処理が常にAPIタイムアウト内で完了し、利用者が同期結果を必須とし、下流容量にも十分な余裕がある低負荷APIなら同期統合でよい。

D 正解

正当なバーストを永続化し、ワーカー処理率をDB限界へ合わせ、再試行・DLQ・監視を独立させる。

7 今後使える判断基準

『入口の公平性』はAPI Gateway REST APIのthrottling/usage plan、『受理後の流量平準化』はSQS、『下流保護』はコンシューマー同時実行で分担する。usage planのクォータはベストエフォートなので、厳密な課金上限・アクセス認可には別の計測と認可を使う。

8 関連サービスとの比較

API Gateway throttling/usage planはREST APIの過負荷・利用抑制であり、厳密な費用統制ではない。SQSは耐久バッファ、Lambda reserved concurrencyやワーカー数は下流への放出率制御、DLQは毒メッセージ隔離を担う。429と202はそれぞれ入口拒否と耐久受理を表し、単一のタイムアウト設定では代替できない。

QUESTION 061

問題61

PERFORMANCE

本番相当

単一選択

計測結果から選ぶメモリ最適化インスタンス

リスク計算アプリはEC2上でJavaヒープ内に90GiBの参照データを保持する。CloudWatch AgentとJVM計測では、ピーク時の常駐メモリが118GiB、CPU使用率は平均32%・最大55%、ネットワークは1Gbps未満、永続データは500GBのgp3 EBSで6,000 IOPSあれば十分である。現在の汎用インスタンスではメモリ不足によるGC停止とスワップが発生する一方、vCPUの多くは未使用である。アプリはGPUを使わず、ローカルインスタンスストアが失われても困らない。同時実行スレッド数を増やしてもCPU使用率は上がらず、GCログではヒープ圧迫が停止時間の主要因であることも確認している。性能問題を解消し、過剰なCPU・ローカルディスク費用を避ける最初のインスタンスファミリー選択はどれか。

- A** C系のコンピューティング最適化を選び、現在より多いvCPUでGCを高速化する。メモリ対vCPU比は下がってもCPUクロックを優先する。
- B** R系などのメモリ最適化を選び、ヒープ、OS、GCの余裕を含むRAM容量ヘライトサイズする。EBS性能はgp3側で必要値を設定する。
- C** I系のストレージ最適化を選び、大容量NVMeインスタンスストアへヒープのswapを置く。再起動時は参照データを再構築する。
- D** P系のアクセラレーテッドコンピューティングを選び、GPUメモリをJavaヒープとして透過的に利用する。アプリコードは変更しない。

ここで答えを決める

正解・解説は次のページから

ANSWER 061

問題61の正解・解説

1 正解 **B**

- B** R系などのメモリ最適化を選び、ヒープ、OS、GCの余裕を含むRAM容量ヘライトサイズする。EBS性能はgp3側で必要値を設定する。

2 問題文の重要な条件

- ボトルネックは118GiBのメモリ使用とGC/スワップで、CPUは余っている
- EBS 6,000 IOPSと1Gbps未満のネットワークで十分である
- GPUや大容量ローカルNVMeを利用するアプリ要件はない

3 正解へ至る判断過程

1. インスタンス名ではなく観測値から支配資源を特定すると、CPU、ネットワーク、ストレージではなくRAM容量が不足している。
2. メモリ最適化ファミリーはvCPU当たりのメモリ比が高く、不要なvCPUを増やさずヒープとOS余裕を確保しやすい。
3. 永続I/O要件はEBS gp3で独立に設定できるため、ローカルストレージ最適化へ移る必要はない。

4 正解が要件を満たす理由

Bは問題となっているメモリ容量と帯域に資源配分を寄せ、スワップを避けられるRAMを確保しながら、CPU中心ファミリーの不要なvCPUを減らせる。ヒープ上限だけでなくOS、ページキャッシュ、ネイティブメモリ、GCの安全余裕を含めてサイズを選ぶ。既に十分と判明した永続I/Oはgp3の容量・IOPS・スループットを独立調整できるため、メモリとブロック性能を適材適所で設定できる。

5・6 各選択肢の検討

不適切な理由と、正解になり得る条件を対にして確認します。

A 不正解

C系はCPU集約処理に適するが、本件はCPUが余りメモリ不足である。vCPUを増やしてもヒープ容量不足とスワップの根本解決にならない。

この選択肢が正解になる条件

CPUが継続的に高使用率で、メモリ使用量が小さいエンコード、科学計算、バッチ処理ならC系が適する。

B 正解

実測された支配資源であるRAMへ高い比率で投資し、EBS性能を別軸で設定して過剰CPUとローカルディスクを避ける。

C 不正解

I系は低遅延・高IOPSのローカルデータセット向けで、swap依存はRAM不足を隠すだけでGC停止を解消しない。永続性のないNVMeも不要である。

この選択肢が正解になる条件

非常に大きいローカルNoSQL、検索インデックス、データウェアハウスの一時領域など、低遅延ランダムI/Oが支配的ならI系が有力である。

D 不正解

GPUメモリは通常のJavaヒープへ透過的に追加されず、GPU対応コードもない。高価なアクセラレータが未使用になる。

この選択肢が正解になる条件

CUDA等に対応した機械学習、レンダリング、科学計算でGPU演算とGPUメモリを明示的に使うならP/G系が適する。

7 今後使える判断基準

EC2ファミリーはCPU、メモリ、ネットワーク、EBS、ローカルI/Oの実測ボトルネックで選ぶ。高CPU=C、メモリ常駐=R、低遅延ローカルI/O=I、GPU/アクセラレータ=P/G、均衡=Mを出発点にする。

8 関連サービスとの比較

C系は高CPU比、R系は高メモリ比、I系はローカルNVMe/IOPS、P/G系はGPU、M系は汎用バランスである。EBS最適化やgp3設定はファミリー選択とは別に帯域上限を確認する。

QUESTION 062

問題62

SECURE

本番より難しい

単一選択

暗号化EBSスナップショットの跨アカウント・跨リージョン移送

アカウントAのap-northeast-1にあるEBSスナップショットはAWSマネージドキーaws/ebsで暗号化されている。災害対策として、アカウントBのap-southeast-2に独立コピーを作り、Bが所有するカスターマネージドKMSキーで再暗号化したい。Bの復旧担当だけがボリュームを作成でき、Aの管理者が侵害されてもB側コピーを削除できない境界にする。現在、Aが元スナップショットをBへ直接共有しようとして失敗している。中間コピーの作成と共有は監査用ロールだけが実施し、完了後はB側の復旧担当がA側キーを継続利用しなくてよい状態を目指す。暗号化を解除せずに実現する正しい手順はどれか。

- A** aws/ebsのキーポリシーにアカウントBを追加し、元の暗号化スナップショットを公開共有する。Bはap-southeast-2から同じスナップショットARNを直接使用する。
- B** 元スナップショットから非暗号化スナップショットを作成してBへ公開し、B側で暗号化し直す。転送中だけS3のSSE-S3を使う。
- C** Aで元スナップショットをカスターマネージドキーへ再暗号化コピーし、そのスナップショットとキー利用権限をBへ限定共有する。Bがap-southeast-2へコピーし、B所有の同リージョンKMSキーで再暗号化する。
- D** 元スナップショットだけをBへ共有し、KMS権限は与えない。BはEC2サービスリンクロールを使えばaws/ebs暗号化を越えてコピーできる。

ここで答えを決める

正解・解説は次のページから

ANSWER 062

問題62の正解・解説

1 正解 C

- C** Aで元スナップショットをカスターマネージドキーへ再暗号化コピーし、そのスナップショットとキー利用権限をBへ限定共有する。Bがap-southeast-2へコピーし、B所有の同リージョンKMSキーで再暗号化する。

2 問題文の重要な条件

- 元スナップショットは編集不能なAWSマネージドキーaws/ebsで暗号化されている
- 暗号化を維持してアカウントとリージョンの両方をまたぐ
- 最終コピーはアカウントB所有・宛先リージョンのKMSキーで独立保護する

3 正解へ至る判断過程

1. aws/ebsのキーポリシーは顧客が外部アカウント向けに変更できないため、そのままの暗号化スナップショットを共有できない。
2. まずA内で共有可能なカスターマネージドキーへ暗号化コピーし、スナップショット許可とKMS許可の両方をBへ与える。
3. Bは共有スナップショットを自アカウント・宛先リージョンへコピーする際にB所有キーを指定し、以後の削除・復旧権限を分離する。

4 正解が要件を満たす理由

Cは暗号化を一度も解除せず、AWSマネージドキーの共有制約をカスターマネージドキーへのコピーで解消する。共有段階ではAのスナップショット属性にBを追加し、AのKMSキーポリシーとB側IAMの双方で必要なDecrypt/ReEncrypt/CreateGrant等を許可する。Bが宛先リージョンにコピーしてBのKMSキーを指定すれば、コピーの所有権と鍵管理がBへ移り、Aから独立した復旧境界になる。

5・6 各選択肢の検討

不適切な理由と、正解になり得る条件を対にして確認します。

A 不正解

AWSマネージドキーaws/ebsのキーポリシーは編集できず、暗号化スナップショットを公開共有することもできない。スナップショットARNはリージョンもまたがない。

この選択肢が正解になる条件

元が共有可能なカスタマーマネージドキーで暗号化され、同リージョンの特定アカウントだけへ限定共有する場合は、適切な鍵・スナップショット共有が可能である。

B 不正解

暗号化スナップショットのコピーから非暗号化スナップショットを作ることはできず、機密データを一時的に平文化する要件でもない。

この選択肢が正解になる条件

元データが最初から非暗号化で規制上許可され、移行先で初めて暗号化する別ケースなら、コピー時暗号化を検討できる。

C 正解

共有可能キーへの中間コピー、スナップショットとKMSの二重許可、宛先アカウント・リージョンでの再暗号化という正しい順序を満たす。

D 不正解

暗号化スナップショットの利用にはデータキーを復号するKMS権限が必須で、サービスリンクロールが所有者のキーポリシーを迂回することはない。

この選択肢が正解になる条件

スナップショットが非暗号化ならKMS権限は不要だが、本件の暗号化・再暗号化要件では成立しない。

7 今後使える判断基準

暗号化スナップショット共有は『スナップショットの共有権限』と『KMSキーの利用権限』を別々に確認する。aws/ebs等のAWSマネージドキーなら、先にカスタマーマネージドキーへコピーしてから共有・再暗号化する。

8 関連サービスとの比較

AWSマネージドキーは簡便だがキーポリシー編集不可、カスタマーマネージドキーは外部アカウントを許可できる。共有は参照権、コピーは宛先アカウント所有の独立スナップショットを作る操作である。

QUESTION 063

問題63

RESILIENT

複合難問

単一選択

長時間・監査可能な承認ワークフロー

保険会社は保険金支払処理をサーバーレス化する。1件の処理は書類検証、外部調査、担当者承認、支払、通知の順に進み、担当者承認ではtask tokenを発行して最大45日待つ。外部の支払API自体は非幂等であるため、会社は支払IDをDynamoDBの条件付き書き込みで一意に確保し、タイムアウト時には状態照会または補償処理を行える。支払Taskを盲目的に再試行せず、幂等な検証・通知だけを自動再試行したい。監査部門は各状態の入力、出力、開始・終了、失敗と再試行を実行単位で追跡する。年間件数は30万で、大量ストリーミング処理より耐久性と可視性を優先する。承認待ちの間にコンピュータを占有せず、期限超過時は自動エスカレーションする。最も適切な方式はどれか。

- A** Step Functions Express Workflowを非同期実行し、45日間Wait状態を維持する。CloudWatch Logsを有効化すれば5分の実行上限を越えられる。
- B** 1つのSQSメッセージのVisibility Timeoutを45日に設定し、1個のLambda実行を承認完了まで待機させる。メッセージ受信回数を監査履歴とする。
- C** EventBridgeの1つのルールに全状態を持たせ、イベントの順序だけから現在状態を毎回推測する。失敗時の支払API重複は許容する。
- D** Step Functions Standard Workflowを使用し、callback with task tokenで承認を待つ。Retryは幂等な検証・通知に限定し、支払Taskは一意的な支払ID、状態照会、Catchと補償で扱い、実行履歴を監査する。

ここで答えを決める

正解・解説は次のページから

ANSWER 063

問題63の正解・解説

1 正解 **D**

- D** Step Functions Standard Workflowを使用し、callback with task tokenで承認を待つ。Retryは冪等な検証・通知に限定し、支払Taskは一意的な支払ID、状態照会、Catchと補償で扱い、実行履歴を監査する。

2 問題文の重要な条件

- 人手承認をtask tokenで最大45日待つ長時間ワークフローである
- 非冪等な外部支払は盲目的に再試行せず、一意な業務キー、状態照会、補償で重複を防ぐ
- 状態ごとの詳細な実行履歴を監査できる必要がある

3 正解へ至る判断過程

1. 45日はExpress Workflowの短時間実行上限を大きく超え、コールバック待機パターンもStandard向けである。
2. Standard Workflowは最大1年の耐久実行、exactly-once workflow execution、状態遷移履歴と視覚的追跡に適する。
3. task token、Catch、Timeoutを明示し、Retryを冪等Taskだけに限定する。支払は条件付きで確保した業務キーと状態照会・補償で不確定結果を処理する。

4 正解が要件を満たす理由

DはStandard Workflowの長時間実行とcallback with task tokenを使い、Lambda等を45日間占有せず承認イベントまで状態を保持する。各Task、Retry、Catch、入力・出力は実行履歴から監査できる。Standardのexactly-once workflow executionだけで外部副作用までexactly-onceになるわけではないため、支払Taskには自動Retryを付けず、一意な支払IDの条件付き確保、結果照会、Catchと補償を組み合わせる点も要件に合う。

5・6 各選択肢の検討

不適切な理由と、正解になり得る条件を対にして確認します。

A 不正解

Express Workflowは最大5分で、ログを有効にしても実行上限は延長されない。長期コールバックと完全な実行履歴を重視する用途に合わない。

この選択肢が正解になる条件

5分以内に終わる大量のイベント処理で、at-least-onceを許容し、詳細履歴より高い開始レートと低コストを優先するならExpressが適する。

B 不正解

SQSの可視性とLambda実行時間を45日待機には使えず、業務状態や分岐の監査履歴にもならない。再配信で非幂等処理が重複し得る。

この選択肢が正解になる条件

個々の非同期ジョブが短時間で、単純な再試行キューと幂等ワーカーだけで十分ならSQSが適する。

C 不正解

EventBridgeはイベントルーティングには適するが、単一ルールが45日の状態機械、コールバック、exactly-once制御、詳細履歴を自動提供するわけではない。

この選択肢が正解になる条件

状態を持たない多数のイベントを複数ターゲットへ条件ルーティングし、各コンシューマーが独立・幂等に処理する要件ならEventBridgeが適する。

D 正解

長時間コールバック、耐久状態、詳細監査を一体化し、非幂等支払だけは一意キー・照会・補償で外部副作用の重複を制御する。

7 今後使える判断基準

Step Functionsは実行時間、実行セマンティクス、監査、統合パターンで選ぶ。長時間・人手承認・task token・詳細履歴ならStandardを使う。ただしワークフロー実行のexactly-onceと、非幂等な外部APIの副作用は別物であり、Retry対象を選び、業務キー・照会・補償を設計する。

8 関連サービスとの比較

Standardは最大1年・耐久・監査・コールバック、Expressは最大5分・高スループット・at-least-once中心である。SQSは作業バッファ、EventBridgeはルーティングを担う。どのオーケストレーターでも非幂等な外部支払を盲目的に再試行すれば重複し得るため、外部副作用の制御は別途必要である。

QUESTION 064

問題64

PERFORMANCE

本番相当

複数選択 (2つ)

中断可能な大規模並列バッチ

映像会社は毎週50,000本の独立した動画セグメントを変換する。1ジョブは平均45分でCPU集約、全体は36時間以内に終わればよい。入力と出力はS3にあり、処理はセグメント単位で並列化できる。EC2 Spotの中断は許容するが、各ジョブを毎回ゼロからやり直すと期限を超える。運用チームはインスタンスの起動、ジョブ割当、スケールインを自作せず、複数インスタンスタイプへ分散して空き容量を得たい。ジョブ定義はコンテナ化済みで、同じ入力セグメントを複数回処理しても確定済みの出力を壊さない設計へ変更できる。ピーク以外の計算資源は不要である。失敗ジョブだけを再試行し、完了済み作業を再利用してコストを下げるために選ぶべき対策を2つ選択してください。

- A** 最大ピーク分の1台の大規模On-Demand EC2を常時起動し、ローカルcronで50,000ジョブを直列実行する。障害時は最初から全件を再実行する。
- B** AWS BatchのマネージドEC2 Spotコンピュート環境とジョブキューを使用し、複数の適切なインスタンスタイプとSpot配分戦略、ジョブ再試行を設定する。
- C** 変換アプリがセグメント内の進捗または完了マーカを定期的にS3へcheckpointし、再試行時は最後の安全な地点から再開する。処理と出力を冪等にする。
- D** 全ジョブを固定台数のOn-Demandだけで実行し、Spotは一切使わない。コスト削減のためジョブ結果とログを保存せず、失敗判定も行わない。

ここで答えを決める

正解・解説は次のページから

ANSWER 064

問題64の正解・解説

1 正解 **B・C**

- B** AWS BatchのマネージドEC2 Spotコンピュート環境とジョブキューを使用し、複数の適切なインスタンスタイプとSpot配分戦略、ジョブ再試行を設定する。
- C** 変換アプリがセグメント内の進捗または完了マーカを定期的にS3へcheckpointし、再試行時は最後の安全な地点から再開する。処理と出力を冪等にする。

2 問題文の重要な条件

- 50,000個の独立CPUジョブを36時間以内に並列処理する
- Spot中断を許容し、複数タイプへ分散してコストを下げる
- 完了・途中進捗を保存し、失敗分だけ安全に再試行する

3 正解へ至る判断過程

1. ジョブキュー、依存関係、再試行、コンピュートの伸縮を管理サービスへ委ねるためAWS Batchを選ぶ。
2. Spotの供給・中断リスクは複数インスタンスタイプと適切な配分戦略で分散し、ジョブ単位に再試行する。
3. 基盤の再試行だけでは45分の作業を失うため、アプリ側checkpointと冪等な出力で中断コストを抑える。

4 正解が要件を満たす理由

BはAWS Batchがジョブキューから必要なEC2容量を伸縮し、Spotを複数タイプから調達して50,000件を並列実行する。運用チームは個別インスタンスや割当ロジックを管理せず、失敗したジョブだけ再試行できる。CはS3を耐久checkpointと完了台帳にして、中断時に処理済み部分を再利用し、重複起動でも同じ結果へ収束させる。これによりSpotの低価格を活かしながら期限超過リスクを抑えられる。

5・6 各選択肢の検討

不適切な理由と、正解になり得る条件を対にして確認します。

A 不正解

直列処理は36時間の期限と並列性を活かせず、単一インスタンス障害で全体を失う。常時最大EC2も週次バッチには費用効率が悪い。

この選択肢が正解になる条件

ジョブ数が少なく直列性が必須で、常時利用する単一ホスト上の特殊ライセンスが制約なら大規模On-Demandを選ぶ理由になり得る。

B 正解

大量ジョブのキューイング、並列配置、Spot容量の多様化、伸縮、再試行を管理サービスへ委ねる。

C 正解

Spot中断やジョブ再試行で45分を失わず、完了済み出力の重複・破損も冪等性で防ぐ。

D 不正解

On-Demandは中断しにくいだが、明示されたSpotコスト最適化を使わず、結果・ログ・失敗判定を捨てるため信頼性を満たさない。

この選択肢が正解になる条件

絶対期限が厳しく中断を一切許容せず、追加費用を受け入れるジョブならOn-Demandコンピュート環境を優先またはSpotのフォールバックにできる。

7 今後使える判断基準

中断可能な並列バッチは、AWS Batchでキューと伸縮、Spot多様化で容量、checkpointで失われる作業量、冪等性で再試行安全性を設計する。Spotを選ぶだけでは回復設計にならない。

8 関連サービスとの比較

AWS Batchはバッチスケジューリングとコンピュート管理、Spotは低価格と中断リスク、On-Demandは高い確実性、S3 checkpointはインスタンス寿命から進捗を分離する。

QUESTION 065

問題65

SECURE

本番より難しい

単一選択

改ざん不能なクロスアカウントバックアップ

法律事務所はOrganizations内の本番アカウントでRDS、EBS、EFSを使用し、バックアップを7年間保持する。ランサムウェア対策として、本番の管理者資格情報が完全に侵害されても、攻撃者が既存リカバリポイントの削除や保持期間短縮を行えないようにしたい。バックアップ管理はSecurity OUの専用アカウントへ分離し、AWS Backupの一元ポリシーとクロスアカウントコピーを使う。規制承認後の設定変更猶予期間は設けられるが、その終了後は専用アカウントのrootを含め誰もロックを解除できてはならない。バックアップの作成失敗とコピー遅延には中央アラームを設定し、復旧テストでは専用アカウントから隔離環境へリストアできることも定期確認する。最も適切な構成はどれか。

- A** 専用バックアップアカウントのvaultへバックアップコピーを作成し、カスタマーマネージドKMSキーと最小権限のvault access policyを設定する。AWS Backup Vault LockをCompliance modeで猶予期間後に確定し、最小・最大保持期間を強制する。
- B** 本番アカウント内のvaultだけを使用し、Vault LockをGovernance modeにする。すべての管理者へロック変更権限を付与し、緊急時に保持期間を短縮できるようにする。
- C** AWS Backupを使わず、RDSとEBSの自動スナップショット名をS3バケットへ記録し、その一覧ファイルだけにS3 Object Lockを設定する。元スナップショットは本番管理者が削除できるままにする。
- D** 専用アカウントへ日次レポートだけをコピーし、実際のリカバリポイントは本番アカウントに残す。MFA DeleteをIAMポリシーで要求すれば、侵害された管理者による全削除を防げる。

ここで答えを決める

正解・解説は次のページから

ANSWER 065

問題65の正解・解説

1 正解 **A**

- A** 専用バックアップアカウントのvaultへバックアップコピーを作成し、カスタマーマネージド KMSキーと最小権限のvault access policyを設定する。AWS Backup Vault LockをCompliance modeで猶予期間後に確定し、最小・最大保持期間を強制する。

2 問題文の重要な条件

- 本番アカウント侵害からリカバリポイントの所有境界を分離する
- 7年保持中はrootを含めVault Lockを解除・短縮できないCompliance modeが必要である
- 複数バックアップ対象をAWS Backupポリシーとクロスアカウントコピーで一元管理する

3 正解へ至る判断過程

1. 同一アカウント内だけでは侵害された強い管理者がバックアップ制御面にも到達するため、コピー先アカウントを分離する。
2. Governance modeは十分な権限で変更できるが、Compliance modeは猶予期間後にロック設定と保持中データを変更・削除できない。
3. vault access policy、KMSキー、Organizationsのバックアップポリシーを最小権限化し、本番側にはコピー作成権限だけを与える。

4 正解が要件を満たす理由

AはバックアップのコピーをSecurity OUの別アカウント所有にするため、本番アカウントの資格情報だけでは削除できない。Compliance modeのVault Lockは設定した猶予期間終了後に不変となり、保持期間中のリカバリポイントやロックを管理者・root・AWSでも変更できないWORM統制を提供する。KMSキーとvault policyをコピー・復旧担当へ限定し、組織バックアップポリシーで対象と保持を中央適用できる。

5・6 各選択肢の検討

不適切な理由と、正解になり得る条件を対にして確認します。

A 正解

クロスアカウント所有分離とCompliance modeの不変性を組み合わせ、本番侵害と専用アカウント内の過剰管理権限の両方に耐える。

B 不正解

同一アカウントで障害領域を共有し、Governance modeは適切な権限を持つ利用者が変更・削除できるため、完全侵害とroot耐性を満たさない。

この選択肢が正解になる条件

誤操作防止が目的で、権限のある管理者による緊急上書きを意図的に残す非規制バックアップならGovernance modeが適する。

C 不正解

一覧ファイルをロックしても実体のRDS/EBSリカバリポイントは保護されず、削除後に復元できない。サービススナップショットを単なるS3オブジェクトとして扱っていない。

この選択肢が正解になる条件

保護対象自体がS3オブジェクトで、バージョニング済みバケットにCompliance mode Object Lockを設定する要件ならS3 Object Lockが直接適する。

D 不正解

レポートは復旧データではなく、実体の本番アカウントに残る。MFA条件だけでは侵害された管理制御面から独立した不変コピーを作れない。

この選択肢が正解になる条件

単なる削除操作の追加確認と監査が目的で、アカウント完全侵害を脅威モデルに含めない場合はMFAを補助統制として利用できる。

7 今後使える判断基準

ランサムウェア耐性は、バックアップの『別アカウント所有』と『保持中の不変性』を分けて両方満たす。rootでも変更不可が条件ならVault Lock Compliance modeを猶予期間後に確定する。

8 関連サービスとの比較

Vault Lock Compliance modeは確定後変更不可、Governance modeは権限者が上書き可能。クロスアカウントコピーは制御面分離、S3 Object LockはS3オブジェクトのWORM保護であり、AWS Backupの多サービスリカバリポイントとは適用対象が異なる。

QUESTION 066

問題66

RESILIENT

複合難問

単一選択

Kinesisコンシューマー障害からの再生と分離

IoT企業はKinesis Data Streamsへ每秒4MBのセンサーデータを送信し、異常検知、リアルタイム集計、S3アーカイブの3コンシューマーが読む。現在は全コンシューマーが共有スループットでGetRecordsを使い、集計の大量読み取り時に異常検知のIteratorAgeも上昇する。先週、異常検知アプリがデプロイ障害で36時間停止したが、ストリーム保持期間が24時間だったため前半12時間を再生できなかった。今後は最大48時間の停止から復旧し、各コンシューマーの読取負荷を相互に隔離し、KCLのcheckpointから再開したい。書込シャード数は現在の入力に十分である。各コンシューマーは別チームがデプロイし、1つの遅延や再起動がほかの読取レイテンシと取得可能期間を短くしないことも重要である。最適な変更はどれか。

- A** 保持期間は24時間のまま書込シャード数だけを2倍にし、3コンシューマーが同じGetRecords共有スループットを使い続ける。障害中の欠落はCloudWatch Logsから自動復元する。
- B** 保持期間を少なくとも48時間超へ延長し、各アプリを登録済みenhanced fan-out consumerとして構成する。KCL checkpointと冪等処理を使い、復旧時に保持範囲内を再生する。
- C** Kinesis Data Firehose 1本だけに置き換え、3つの低遅延コンシューマーがFirehoseの同じ配信ストリームから任意位置を再読込する。
- D** 保持期間を48時間にする代わりに、各レコードを送信直後に削除してストレージ費を抑える。異常検知停止中はプロデューサーも停止させる。

ここで答えを決める

正解・解説は次のページから

ANSWER 066

問題66の正解・解説

1 正解 **B**

- B** 保持期間を少なくとも48時間超へ延長し、各アプリを登録済みenhanced fan-out consumerとして構成する。KCL checkpointと冪等処理を使い、復旧時に保持範囲内を再生する。

2 問題文の重要な条件

- 最大48時間停止したコンシューマーが欠落なく再生できる保持期間が必要である
- 3コンシューマーの読み取りスループット競合を分離する
- 書込容量は十分で、KCL checkpointからの再開と重複耐性を維持する

3 正解へ至る判断過程

1. 障害時間より保持期間が短いと、コンシューマー実装を改善しても期限切れレコードは再生できないため、まず保持期間をRTO余裕込みで延長する。
2. 共有GetRecordsではコンシューマーがシャードの読取容量を分け合うため、enhanced fan-outで登録コンシューマーごとの専用読取スループットを得る。
3. KCL checkpointを耐久保存し、再開位置から読み直す。再試行やフェイルオーバーで重複し得るので処理は冪等にする。

4 正解が要件を満たす理由

Bはレコードを少なくとも障害許容時間より長く保持するため、36～48時間停止後でもcheckpoint以降をストリームから再生できる。Enhanced fan-outは各登録コンシューマーにシャード当たり専用の読取スループットを与え、集計の読み取りが異常検知を圧迫する共有競合を解消する。KCLはシャード割当てとcheckpointを管理し、再処理時の重複は業務キーによる冪等更新で安全に扱える。

5・6 各選択肢の検討

不適切な理由と、正解になり得る条件を対にして確認します。

A 不正解

シャード追加は容量を増やしても24時間を過ぎたレコードを復活させず、共有読取競合の設計も残る。CloudWatch Logsは自動的なストリーム正本ではない。

この選択肢が正解になる条件

保持24時間以内に必ず復旧でき、主な問題が書込またはシャード当たり処理容量不足である場合はreshardingが適切になる。

B 正解

保持期間で再生可能窓を、enhanced fan-outで読取性能の障害分離を、KCL/冪等性で安全な再開を満たす。

C 不正解

FirehoseはS3等への管理された配信に適するが、複数アプリが任意checkpointから低遅延に再生するストリームコンシューマーモデルを代替しない。

この選択肢が正解になる条件

目的が変換・バッファ後にS3、Redshift、OpenSearch等へ配信することで、独自の複数リアルタイムコンシューマーや再生が不要ならFirehoseが適する。

D 不正解

Kinesisレコードを送信者が個別削除するモデルではなく、プロデューサー停止は他コンシューマーとデータ収集を巻き込み、障害分離に反する。

この選択肢が正解になる条件

全システムが同一の計画停止を受け入れ、データ源側で安全に再送できる特殊なケースでも、通常は保持期間設計を優先する。

7 今後使える判断基準

Kinesis復旧では、保持期間を最大停止時間+復旧余裕以上にし、読取競合にはenhanced fan-out、再開位置にはcheckpoint、再配信には冪等性を使う。シャード追加だけでは時間窓を延ばせない。

8 関連サービスとの比較

共有fan-outは低コストだが読取容量を共有し、enhanced fan-outはコンシューマー専用帯域を提供する。保持期間は再生可能時間、KCLは処理協調、Firehoseは宛先への管理配信を担う。

QUESTION 067

問題67

SECURE

本番相当

単一選択

Organization trailの中央保護

企業はOrganizations配下の80アカウントについて、全リージョンの管理イベントと重要S3バケットのデータイベントを監査する。現在は各アカウント管理者がローカルtrailを作るため、停止・削除や設定漏れが発生し、新規アカウントの追加も手作業である。今後はLog ArchiveアカウントのS3バケットへ一元保存し、メンバーアカウント管理者がtrailやログを変更できないようにする。CloudTrailからの正規配信だけを許可し、改ざん・削除を検出でき、監査担当だけが読み取れる構成にする。ログは7年間保持し、S3バケットの管理権限をワークロードアカウントへ一切付与せず、削除保護と鍵管理もセキュリティ担当へ分離する。運用負荷を最も低く要件を満たす方法はどれか。

- A** 各メンバーアカウントに単一リージョンtrailを個別作成し、ログは同じアカウントのS3へ保存する。管理者にはCloudTrailとS3の全権限を残す。
- B** CloudTrailのEvent historyだけを90日ごとにCSV出力して管理アカウントのEBSへ保存する。データイベントもEvent historyに自動で含まれる前提とする。
- C** 管理アカウントまたはCloudTrail委任管理者からmulti-Region organization trailを作り、Log ArchiveアカウントのS3へ配信する。bucket/KMS policyをtrail ARNと組織プレフィックスへ限定し、log file validationと保護された保持設定を有効にする。
- D** AWS Config aggregatorだけを全アカウントで有効化し、設定変更履歴をCloudTrail API 監査ログの代替にする。S3 GetObjectなどのデータイベントはVPC Flow Logsから復元する。

ここで答えを決める

正解・解説は次のページから

ANSWER 067

問題67の正解・解説

1 正解 C

- C** 管理アカウントまたはCloudTrail委任管理者からmulti-Region organization trailを作り、Log ArchiveアカウントのS3へ配信する。bucket/KMS policyをtrail ARNと組織プレフィックスへ限定し、log file validationと保護された保持設定を有効にする。

2 問題文の重要な条件

- 80アカウントと新規参加アカウントへ一貫した全リージョン監査を自動適用する
- メンバーアカウント管理者からtrail管理とログ保管を分離する
- 中央S3への配信を最小権限化し、ログの改ざん・削除を検出または防止する

3 正解へ至る判断過程

1. Organizations全体へ自動適用し、メンバーが変更できない管理単位としてorganization trailを選ぶ。
2. multi-Region設定と選択したデータイベントを中央Log Archiveバケットへ配信し、アカウント追加時もサービスリンクロールで収集対象にする。
3. S3/KMSポリシーをCloudTrailのSourceArn、組織用プレフィックス、必要操作へ限定し、log file validationと保持保護で証拠性を高める。

4 正解が要件を満たす理由

Cはorganization trailを全メンバーへ適用し、新規アカウントも自動的に記録対象へ加える。通常のメンバーアカウントはorganization trailを停止・変更・削除できない。ログ実体を専用Log Archiveアカウントが所有し、CloudTrailサービスの書込と監査担当の読取だけをbucket/KMS policyで許可すれば職務分離できる。Log file validationは配信後の変更・削除を検出し、S3のバージョニングやObject Lock等の保持統制を組み合わせれば削除耐性も強化できる。

5・6 各選択肢の検討

不適切な理由と、正解になり得る条件を対にして確認します。

A 不正解

アカウントごとのtrailは設定ドリフトと新規追加作業を残し、同じ管理者がtrailとログを停止・削除できる。単一リージョンでは他リージョンも漏れる。

この選択肢が正解になる条件

Organizationsを使用しない独立した1アカウントだけが対象で、同一アカウント管理者による保管を許容する小規模環境ならアカウントtrailで足りる。

B 不正解

Event historyは永続trailの代替ではなく、デフォルトでデータイベントを含まない。手動CSVとEBSは完全性、長期保持、アカウント追加自動化を満たさない。

この選択肢が正解になる条件

直近90日の管理イベントを短時間の運用調査で確認するだけならEvent historyを追加trailなしで利用できる。

C 正解

organization trail、自動アカウント適用、中央所有、最小権限配信、整合性検証を組み合わせ、収集漏れとメンバーによる改変を抑える。

D 不正解

Configはリソース構成と変更履歴を記録するが、API呼出主体・時刻・要求などCloudTrailイベントを代替しない。Flow LogsもS3オブジェクトAPIを再現しない。

この選択肢が正解になる条件

目的がリソースの構成状態、準拠評価、設定履歴の集約であり、API監査とデータイベントが別途CloudTrailで確保される場合はConfig aggregatorが適する。

7 今後使える判断基準

マルチアカウント監査はorganization trailで収集統制を中央化し、ログ保管アカウントを分離する。証拠性は配信元限定のbucket/KMS policy、log file validation、バージョニング／Object Lock等の保持統制に分解する。

8 関連サービスとの比較

Organization trailは全アカウントのAPI監査、CloudTrail Event historyは直近の管理イベント閲覧、AWS Config aggregatorは構成状態集約、VPC Flow Logsはネットワークフロー記録であり、互いに代替ではない。

QUESTION 068

問題68

SECURE

本番より難しい

複数選択 (2つ)

未暗号化RDSの段階的暗号化移行

ある金融会社は、単一AZで稼働するAmazon RDS for MySQLの既存DBインスタンスを監査したところ、保存時の暗号化が無効であることを発見しました。DBは2TBで、アプリケーションはDBエンドポイントを設定ファイルから参照しています。監査期限までに、会社管理のKMSキーでDB、スナップショット、自動バックアップを暗号化する必要があります。エンジン変更やアプリケーションの大幅な改修は認められません。週末に12時間の書き込み停止を確保でき、事前演習ではスナップショットのコピー、復元、検証が8時間以内に完了しました。スナップショット取得から切替判断までは元DBをread-onlyに保ち、検証不合格なら元DBへ戻し、合格後に新DBで書き込みを再開します。平文でのエクスポートは避けます。既存DBに暗号化を直接有効化できないことを前提に、実行すべきものを2つ選択してください。

- A** 既存DBインスタンスを変更し、保存時の暗号化と対象KMSキーを指定する。変更を即時適用し、完了後にMulti-AZへ変更して新しい暗号化済みスタンバイフェイルオーバーする。
- B** 書き込みを停止して既存DBの手動スナップショットを作成し、そのスナップショットを同一リージョンでコピーするときに会社管理のKMSキーを指定して暗号化済みコピーを作成する。
- C** 未暗号化DBから同一エンジンの暗号化済みリードレプリカを直接作成する。レプリカ遅延がゼロになった時点でレプリカを昇格し、元DBを保持したまま接続先を切り替える。
- D** 暗号化済みスナップショットコピーから新しいRDS DBインスタンスを復元する。検証後にアプリケーションのDBエンドポイントを新DBへ切り替え、元DBはロールバック期間中保持する。

ここで答えを決める

正解・解説は次のページから

ANSWER 068

問題68の正解・解説

1 正解 **B・D**

- B** 書き込みを停止して既存DBの手動スナップショットを作成し、そのスナップショットを同一リージョンでコピーするときに会社管理のKMSキーを指定して暗号化済みコピーを作成する。
- D** 暗号化済みスナップショットコピーから新しいRDS DBインスタンスを復元する。検証後にアプリケーションのDBエンドポイントを新DBへ切り替え、元DBはロールバック期間中保持する。

2 問題文の重要な条件

- 既存の未暗号化RDS DBインスタンスはインプレースで暗号化できない
- 会社管理のKMSキーを使用し、平文エクスポートを避ける
- 演習済みの12時間停止枠で、切替判断まで元DBをread-onlyのロールバック先として保持する

3 正解へ至る判断過程

1. 暗号化を直接有効化する案と、暗号化属性の異なるリードレプリカを直接作る案を除外する。
2. RDSスナップショットはコピー時に暗号化できるため、未暗号化データをサービス外へ出さずに暗号化境界を作れる。
3. 暗号化済みコピーから別DBとして復元し、接続先を切り替えれば、元DBを残したまま移行できる。

4 正解が要件を満たす理由

Bで書き込み整合性を確保したスナップショットを取得し、コピー処理で指定KMSキーによる暗号化済みスナップショットを作成できます。Dでそのコピーから新規DBを復元すると、新DBのストレージ、以後の自動バックアップ、スナップショットが暗号化されます。コピーと復元は2TBでも所要時間が固定保証されませんが、本問では事前演習が12時間枠内の完了を示し、切替判断まで元DBをread-onlyにするため差分も生じません。

5・6 各選択肢の検討

不適切な理由と、正解になり得る条件を対にして確認します。

A 不正解

RDSは既存の未暗号化DBインスタンスに対して保存時暗号化を直接有効化できません。Multi-AZ化も可用性を高める機能であり、暗号化属性を変更する手段ではありません。

この選択肢が正解になる条件

対象が新規DBの作成時である、または既に暗号化済みDBでKMSキー変更を伴わないMulti-AZ可用性向上だけが要件であれば適切です。

B 正解

未暗号化スナップショットをコピーする際にKMSキーを指定することで、暗号化済みの復元元をサービス内で作れます。

C 不正解

暗号化されていないRDS DBから、暗号化属性だけを変えたリードレプリカを直接作成する手順にはできません。

この選択肢が正解になる条件

元DBも暗号化済みで、暗号化属性を変えずに読み取り負荷分散または短時間の昇格移行を行う要件なら候補になります。

D 正解

暗号化済みスナップショットからの復元は暗号化済みの別DBを作り、元DBを残した接続先切り替えを可能にします。

7 今後使える判断基準

既存リソースで暗号化をインプレース変更できない場合は、整合性スナップショットを暗号化コピーし、別リソースへ復元して切り替える。大容量ではコピー・復元時間を固定値と見なさず演習し、取得後の変更を止めるかCDCで追跡させ、RTOとデータ整合性を別々に検証する。

8 関連サービスとの比較

RDSのMulti-AZは同期スタンバイによる可用性向上、リードレプリカは非同期複製による読み取り拡張が主目的です。どちらも未暗号化DBの暗号化属性を直接変更する機能ではなく、スナップショットの暗号化コピーと復元が適合します。

QUESTION 069

問題69

RESILIENT

複合難問

単一選択

RTOと常時費用で選ぶリージョンDR

オンライン卸売会社は、3層Webシステムを単一リージョンの複数AZで運用しています。アプリケーション層はEC2 Auto Scaling、データ層はAmazon Aurora MySQLです。リージョン全体の障害に備え、経営層はRTO 15分、RPO 5分を承認しました。災害時にも主要な注文APIと参照画面を同じ機能で提供する必要があり、復旧時にOSやミドルウェアを一から構築する時間はありません。一方、平常時に本番と同規模のコンピューティング能力を二重稼働させる予算はなく、DR環境は毎月の訓練で実際にリクエストを処理して検証したいと考えています。データは継続的にDRリージョンへ複製でき、DNS切り替えとスケールアウトは自動化できます。これらの条件を最も費用対効果高く満たすDR戦略はどれですか。

- A** バックアップと復元を採用し、AuroraのスナップショットとEC2 AMIをDRリージョンへ毎日コピーする。障害宣言後にCloudFormationで全リソースを作成してデータを復元する。
- B** パイロットライトを採用し、データベース複製と最小限の中核サービスだけを常時維持する。障害宣言後にアプリケーション層とロードバランサーを新規作成して再構成する方式とする。
- C** マルチサイトのアクティブ・アクティブを採用し、両リージョンに本番規模の全層を常時稼働させ、Route 53で通常時から利用者トラフィックを両方の環境へ振り分ける。
- D** ウォームスタンバイを採用し、DRリージョンに縮小済みだが完全に機能する全層と継続レプリケーションを常時稼働させる。障害時にスケールアウトしてDNSを切り替える。

ここで答えを決める

正解・解説は次のページから

ANSWER 069

問題69の正解・解説

1 正解 **D**

- D** ウォームスタンバイを採用し、DRリージョンに縮小済みだが完全に機能する全層と継続レプリケーションを常時稼働させる。障害時にスケールアウトしてDNSを切り替える。

2 問題文の重要な条件

- RTO 15分、RPO 5分という短い復旧目標
- DR側で平常時から機能確認でき、構築作業を復旧経路に入れない
- 本番規模の二重稼働は不可だが、縮小環境の常時費用は許容される

3 正解へ至る判断過程

1. バックアップ復元は低コストだが、プロビジョニングと復元により15分RTOを満たしにくい。
2. パイロットライトは中核データを保持するが、アプリケーション全層を障害後に構築するため月次の実処理検証と短いRTOに不利である。
3. 縮小した完全機能環境を常時動かすウォームスタンバイなら、スケールと経路変更だけを復旧手順にできる。

4 正解が要件を満たす理由

ウォームスタンバイは、DRリージョンに縮小版ながらエンドツーエンドで動作する環境を維持します。継続的なデータ複製により5分のRPOを狙え、障害時は事前構築済み環境の容量拡大とDNS切り替えに限定できるため15分のRTOに適合します。通常時にテストトラフィックを流せる一方、本番同規模のアクティブ・アクティブより常時費用を抑えられます。

5・6 各選択肢の検討

不適切な理由と、正解になり得る条件を対にして確認します。

A 不正解

日次コピーではRPO 5分を満たさず、インフラ作成と2TB級データ復元がRTO 15分を超える可能性が高いです。

この選択肢が正解になる条件

RTOが数時間から1日、RPOが24時間程度で、平常時コストを最小にすることが最優先なら正解になります。

B 不正解

データ層は準備できますが、全アプリケーション層の作成と構成が障害後に残り、15分RTOと通常時の全機能検証に不利です。

この選択肢が正解になる条件

RTOが数十分から数時間で、データ中核だけを常時維持し、障害後の自動プロビジョニング時間を許容できるなら適切です。

C 不正解

RTO/RPOは満たしやすいものの、本番規模を両リージョンで常時稼働するため、明示された予算制約を満たしません。

この選択肢が正解になる条件

ほぼゼロのRTO/RPOが必要で、両リージョンを通常時から利用し、二重の本番容量コストを許容できるなら最適です。

D 正解

完全機能の縮小環境を常時維持するため短いRTOと訓練可能性を満たし、フルサイズ二重稼働より費用を抑えます。

7 今後使える判断基準

DR戦略は名称ではなく、RPOを決めるデータ複製方式、RTO内に残る構築作業、平常時の稼働規模の3点で選ぶ。短いRTOでフル二重費用が不可ならウォームスタンバイが有力になる。

8 関連サービスとの比較

バックアップと復元は最安だがRTO/RPOが長く、パイロットライトはデータ中核のみ常時稼働、ウォームスタンバイは縮小した完全機能環境、アクティブ・アクティブは全容量を常時提供します。復旧目標が短くなるほど一般に費用と運用複雑性が上がります。

QUESTION 070

問題70

PERFORMANCE

本番相当

単一選択

世界各地からの大容量S3アップロード

映像制作会社は東京リージョンのAmazon S3バケットへ、世界20都市の撮影現場から素材をアップロードしています。1ファイルは20～200GBで、現場のインターネット回線は十分な帯域があるものの、長距離通信の遅延と一時的な切断により単一PUTの再送が頻発しています。利用者は既存のデスクトップアプリケーションから直接S3へ送信しており、各都市にVPN、Direct Connect、プロキシサーバーを構築する運用は避けたいと考えています。バケットのリージョンは変更できず、転送完了後のオブジェクトキーも維持する必要があります。失敗時はファイル全体ではなく未完了部分だけを再送し、複数の部分を並列転送して帯域を使い切ることで、世界各地からAWSの最適化されたネットワークへ早く入れることが要求されます。最小の運用負荷で要件を満たす構成はどれですか。

- A** S3 Transfer Accelerationをバケットで有効化し、クライアントをアクセラレーションエンドポイントへ変更する。S3マルチパートアップロードでパートを並列送信し、失敗したパートだけを再送する。
- B** 各都市にAWS DataSyncエージェントを常設し、ローカルNFSにファイルを集約してから、公共インターネット経由のDataSyncタスクで東京のS3へ毎時間転送する。
- C** Amazon CloudFrontディストリビューションのオリジンにS3を設定し、撮影アプリケーションからCloudFrontの通常のビューワーエンドポイントへPUTしてエッジキャッシュに保存する。
- D** AWS Global Acceleratorを作成し、東京リージョンのS3バケットを直接エンドポイントとして登録する。単一PUTを維持したまま静的Anycast IPへアップロードする。

ここで答えを決める

正解・解説は次のページから

ANSWER 070

問題70の正解・解説

1 正解 **A**

- A** S3 Transfer Accelerationをバケットで有効化し、クライアントをアクセラレーションエンドポイントへ変更する。S3マルチパートアップロードでパートを並列送信し、失敗したパートだけを再送する。

2 問題文の重要な条件

- 20～200GBの単一大容量オブジェクトで、切断時の部分再送が必要
- 世界各地から単一リージョンのS3へ長距離転送する
- 拠点ごとのインフラを置かず、既存クライアントの小変更で対応する

3 正解へ至る判断過程

1. 部分再送と並列化はS3マルチパートアップロードが直接満たす。
2. 地理的距離による転送遅延には、エッジロケーションからAWSネットワークを利用するS3 Transfer Accelerationが合う。
3. DataSyncはファイルシステム移行には有効だが、各拠点エージェントが不要という条件に反し、CloudFrontやGlobal AcceleratorはこのS3直接アップロード要件の代替ではない。

4 正解が要件を満たす理由

S3マルチパートアップロードは各パートを独立・並列に送信でき、ネットワーク障害時に失敗パートのみを再送できます。Transfer Accelerationは世界各地のクライアントから近いエッジロケーションを利用し、その後を最適化されたAWSネットワークでS3バケットへ転送します。バケットの移動や拠点設備が不要で、クライアントのエンドポイントとアップロード方式の変更だけで要件を満たします。

5・6 各選択肢の検討

不適切な理由と、正解になり得る条件を対にして確認します。

A 正解

マルチパートが並列化と部分再送、Transfer Accelerationが長距離経路最適化を担当し、条件を分担して満たします。

B 不正解

DataSyncは反復転送と整合性検証に有効ですが、20都市へのエージェント設置とNFS運用が必要になり、直接アップロードと最小運用負荷の条件に反します。

この選択肢が正解になる条件

各拠点に既存NFSがあり、多数ファイルを定期同期し、帯域制御や転送検証を中央管理する移行要件なら適切です。

C 不正解

CloudFrontの通常用途はコンテンツ配信とキャッシュであり、巨大なS3オブジェクトをビューワーPUTでエッジキャッシュへ蓄積する設計は本要件に合いません。

この選択肢が正解になる条件

S3上の完成済み映像を世界の視聴者へ低遅延配信し、GETのオリジン負荷を減らす要件ならCloudFrontが正解です。

D 不正解

Global Acceleratorの標準エンドポイントとしてS3バケットを直接登録する構成ではなく、単一PUTの部分再送問題も解決しません。

この選択肢が正解になる条件

対象がALB、NLB、EC2、Elastic IPなどで提供する非HTTPまたは動的アプリケーションで、静的Anycast IPとグローバル経路最適化が必要なら候補です。

7 今後使える判断基準

大容量オブジェクトでは、転送単位の並列化・再開性と、地理的なネットワーク経路を別々に判断する。S3への長距離直接転送ならマルチパートとTransfer Accelerationの組み合わせを検討する。

8 関連サービスとの比較

S3 Transfer AccelerationはS3への長距離転送、CloudFrontは主にキャッシュ可能な配信、Global Acceleratorは対応するリージョナルアプリケーションエンドポイントへの経路最適化、DataSyncはストレージ間の管理された反復転送に向きます。

QUESTION 071

問題71

COST

本番より難しい

単一選択

CloudFrontによるオリジン費用削減

ニュース会社はALB配下のEC2で公開サイトを運用しています。アクセスログを分析すると、転送量の75%はCSS、JavaScript、画像など、更新が1日1回以下の公開静的コンテンツです。しかし、すべてのリクエストに端末識別用Cookieと複数の分析用クエリ文字列が付き、ALBとEC2へのGET、オリジンからのデータ転送、ピーク時のスケールアウトが費用の中心になっています。静的コンテンツの内容はCookieで変わらず、バージョン付きファイル名を採用できるため、更新時に長いTTLの失効を待つ必要はありません。動的な記事APIとログイン画面はキャッシュしてはいけません。TLSは維持し、利用者側のアプリ変更も避けます。世界各地の応答時間を改善しつつ、オリジン負荷と総データ転送費を最も効果的に削減する構成はどれですか。

- A** AWS Global Acceleratorの背後にALBを登録し、静的コンテンツと動的APIを同じリスナーで配信する。EC2 Auto Scalingの最小容量は現状のまま維持する。
- B** CloudFrontをALBの前段に置き、静的パスではCookieと不要なクエリ文字列をキャッシュキーおよびオリジン転送から除外し、長いTTLを設定する。動的パスは別ビヘイビアでキャッシュを無効化する。
- C** 静的コンテンツを各リージョンのS3バケットへCRRで複製し、Route 53レイテンシールーティングで利用者を最寄りバケットのS3ウェブサイトエンドポイントへ送る。
- D** ALBのクロスゾーン負荷分散を無効にし、各AZのEC2台数を増やす。静的ファイルには短いCache-Controlを設定し、利用者のブラウザから毎回更新確認を行う。

ここで答えを決める

正解・解説は次のページから

ANSWER 071

問題71の正解・解説

1 正解 **B**

- B** CloudFrontをALBの前段に置き、静的パスではCookieと不要なクエリ文字列をキャッシュキーおよびオリジン転送から除外し、長いTTLを設定する。動的パスは別ビヘイビアでキャッシュを無効化する。

2 問題文の重要な条件

- 転送量の大半がCookieや分析クエリに依存しない静的コンテンツ
- バージョン付きファイル名を採用でき、長いTTLを安全に使える
- 動的APIとログイン画面はキャッシュ不可

3 正解へ至る判断過程

1. 静的・動的パスを分離し、静的部分だけをエッジキャッシュする必要がある。
2. 不要なCookieやクエリをキャッシュキーに含めると同一オブジェクトが細分化されるため、除外してヒット率を上げる。
3. CloudFrontのヒットはALB/EC2へのリクエストとオリジン転送を減らし、配信遅延とスケール費用を同時に改善する。

4 正解が要件を満たす理由

Bはパス別ビヘイビアでキャッシュ可能性を正しく分離し、静的パスのキャッシュキーを実際に内容を変える要素だけに限定します。長いTTLとバージョン付きファイル名により高いキャッシュヒット率を維持でき、オリジンへのリクエスト、ALB/EC2負荷、オリジン側の転送量を削減します。動的パスはキャッシュ無効のため、個人化や認証の正しさも保てます。

5・6 各選択肢の検討

不適切な理由と、正解になり得る条件を対にして確認します。

A 不正解

Global Acceleratorはネットワーク経路と静的IPを提供しますがコンテンツをキャッシュしないため、同じGETが引き続きALBとEC2へ到達し、オリジン負荷削減になりません。

この選択肢が正解になる条件

TCP/UDPまたはキャッシュ不可能な動的通信で、固定Anycast IP、迅速なリージョンフェイルオーバー、経路最適化が主目的なら適切です。

B 正解

キャッシュキー最小化、長いTTL、パス分離により、正確性を保ちながらヒット率と費用削減を最大化します。

C 不正解

複数バケットの保存・複製費用とDNS運用が増え、全世界のエッジキャッシュになりません。さらにS3ウェブサイトエンドポイント単独では要件に応じたHTTPS設計にも制約があります。

この選択肢が正解になる条件

リージョンごとに独立したデータ主権要件があり、複製済みオブジェクトへ地域別に直接到達させることが目的なら候補になります。

D 不正解

EC2増設と短いTTLはオリジンへの条件付きリクエストを増やし、課題である負荷と費用を改善しません。

この選択肢が正解になる条件

コンテンツがすべて秒単位で変化しキャッシュ不能で、AZごとの容量分離が性能上必要という別要件なら検討できます。

7 今後使える判断基準

CDNの費用効果は導入有無だけでなく、内容を変えないCookie・ヘッダー・クエリをキャッシュキーから外し、キャッシュ可能なパスと不可のパスを分離できるかで決まる。

8 関連サービスとの比較

CloudFrontはエッジキャッシュでオリジン要求を削減します。Global Acceleratorはキャッシュせずネットワーク経路を最適化します。S3 CRRはデータを別リージョンへ複製しますが、CDNの多数エッジによるリクエスト集約とは目的と費用構造が異なります。

QUESTION 072

問題72

SECURE

複合難問

複数選択 (2つ)

ワークロードの長期アクセスキー撤廃

医療SaaS企業の監査で、EC2上のJavaサービスと複数のLambda関数が、同じIAMユーザーのアクセスキーを環境変数と設定ファイルに保存してS3およびSQSへアクセスしていることが判明しました。キーは2年間ローテーションされておらず、漏えい時には全ワークロードを停止して設定を書き換える必要があります。会社は長期認証情報をコードと設定から完全に除去し、各ワークロードに必要なリソースと操作だけを許可し、認証情報を自動更新したいと考えています。EC2アプリケーションとLambdaのコードはAWS SDKの標準認証情報プロバイダーチェーンを使用するよう変更できます。複数サービスで同じ権限を共有する必要はなく、人間によるAWSアクセスはこの設問の対象外です。最小権限と運用負荷の要件を満たす変更を2つ選択してください。

- A** 既存IAMユーザーのアクセスキーをSecrets Managerへ移し、EC2とLambdaに同じシークレットの読み取りを許可する。Secrets ManagerのLambdaローテーションで90日ごとにキーを再発行する。
- B** EC2とLambdaごとに別のIAMユーザーを作り、アクセスキーを暗号化した環境変数として配布する。IAM Access Analyzerで未使用権限を四半期ごとに削除する。
- C** EC2サービス専用の最小権限IAMロールを作り、インスタンスプロファイルとしてEC2に関連付ける。SDKがインスタンスメタデータから自動更新される一時認証情報を取得するようにする。
- D** Lambda関数群を用途別に分け、各関数または同一権限群に最小権限の実行ロールを関連付ける。SDKには固定キーを渡さず、実行環境が提供する一時認証情報を使用する。

ここで答えを決める

正解・解説は次のページから

ANSWER 072

問題72の正解・解説

1 正解 C・D

- C** EC2サービス専用の最小権限IAMロールを作り、インスタンスプロファイルとしてEC2に関連付ける。SDKがインスタンスメタデータから自動更新される一時認証情報を取得するようにする。
- D** Lambda関数群を用途別に分け、各関数または同一権限群に最小権限の実行ロールを関連付ける。SDKには固定キーを渡さず、実行環境が提供する一時認証情報を使用する。

2 問題文の重要な条件

- 長期アクセスキーをコードと設定から完全に除去する
- EC2とLambdaで権限境界が異なり、最小権限が必要
- AWS SDKの標準認証情報プロバイダーチェーンへ変更可能

3 正解へ至る判断過程

1. アクセスキーを別の保存場所へ移すだけでは、長期認証情報のライフサイクルと漏えいリスクが残る。
2. AWSコンピューティングサービスには、サービスが引き受けて一時認証情報を供給するIAMロールを割り当てられる。
3. EC2はインスタンスプロファイル、Lambdaは実行ロールを使い、用途別ポリシーに分ければ自動更新と最小権限を両立できる。

4 正解が要件を満たす理由

CではEC2のインスタンスプロファイルを通じてロールの一時認証情報が提供され、DではLambda実行ロールの一時認証情報が実行環境に提供されます。標準プロバイダーチェーンを使えばアプリケーションがキー値やローテーションを管理する必要がありません。さらにEC2サービスとLambda関数群に別ロールを設けるため、一方の侵害で他方の権限まで得る共有IAMユーザーの問題を解消できます。

5・6 各選択肢の検討

不適切な理由と、正解になり得る条件を対にして確認します。

A 不正解

保管場所とローテーションを改善しても、共有された長期アクセスキー自体が残り、読み取り権限を持つ全ワークロードへ影響範囲が広がります。

この選択肢が正解になる条件

外部SaaSなどIAMロールを利用できないクライアントのAPI資格情報を安全に保存し、自動ローテーションする要件ならSecrets Managerが適切です。

B 不正解

IAMユーザーを分けても長期キーの配布と更新が残り、環境変数から漏えいする可能性と運用負荷を解消できません。

この選択肢が正解になる条件

AWS外部のレガシー製品が一時認証情報を利用できず、短期的な移行措置として厳格なローテーションと個別ユーザーが必要なら限定的に成立します。

C 正解

EC2用IAMロールとインスタンスプロファイルにより、SDKが自動更新される一時認証情報を利用できます。

D 正解

Lambda実行ロールを権限単位で分離すれば、キー配布なしで関数ごとの最小権限を実装できます。

7 今後使える判断基準

AWS上のワークロードには長期アクセスキーを配らず、実行主体にネイティブなIAMロールを付与し、SDKの標準認証情報チェーンで一時認証情報を取得させる。ロールは信頼境界と権限集合ごとに分ける。

8 関連サービスとの比較

IAMロールとSTS一時認証情報は自動失効・更新され、EC2インスタンスプロファイルやLambda実行ロールから利用できます。Secrets ManagerはDBパスワードや外部APIキーには有効ですが、利用可能なIAMロールの代わりにAWSアクセスキーを保管する用途は優先されません。

QUESTION 073

問題73

RESILIENT

本番相当

単一選択

DynamoDBによるマルチリージョン書き込み

ゲーム会社は北米と欧州の2リージョンでAPIを稼働させ、プレイヤーの設定と進行状況をDynamoDBに保存しています。現在は北米だけが書き込みを受け、欧州からは往復遅延が大きく、北米リージョン障害時には手動復旧に40分かかります。今後は両リージョンで通常時からローカルの読み取りと書き込みを受け付け、片方のリージョンが停止してもDNSのヘルスチェックにより数分以内にもう一方だけで継続したいと考えています。通常時の利用者は最寄りリージョンへ誘導します。アプリケーションは同一項目への同時更新をバージョン番号で検出でき、リージョン間複製の短い遅延と競合解決を許容します。DBサーバーや独自レプリケーション基盤は運用したくなく、既存のキー値アクセスも維持します。最も運用負荷が低く要件を満たす構成はどれですか。

- A** 北米のDynamoDBテーブルをPoint-in-Time Recoveryで保護し、障害時に欧州へ復元する。Route 53フェイルオーバーレコードが復元完了後に欧州APIへ切り替わるようにする。
- B** 北米テーブルのDynamoDB StreamsをKinesis Data Firehoseで欧州S3へ配信し、障害時にS3 Importで新しい欧州テーブルを作成してAPIの接続先を変更する。
- C** 両リージョンをレプリカとするDynamoDB Global Tablesを構成し、各APIをローカルレプリカへ接続する。Route 53で通常時はレイテンシー、障害時は正常リージョンヘルレーティングする。
- D** Aurora Global Databaseへデータを移行し、北米をライター、欧州を読み取り専用セカンダリにする。北米障害時は欧州を昇格して新しい書き込み先に切り替える。

ここで答えを決める

正解・解説は次のページから

ANSWER 073

問題73の正解・解説

1 正解 C

- C** 両リージョンをレプリカとするDynamoDB Global Tablesを構成し、各APIをローカルレプリカへ接続する。Route 53で通常時はレイテンシー、障害時は正常リージョンヘルレーティングする。

2 問題文の重要な条件

- 両リージョンで通常時からローカル書き込みを受け付けるactive-active要件
- 短い非同期複製遅延と競合解決をアプリケーションが許容する
- 独自レプリケーションやDBサーバー運用を避ける

3 正解へ至る判断過程

1. バックアップ復元やS3インポートは障害後の再構築であり、通常時の両リージョン書き込みと数分の復旧を満たさない。
2. Aurora Global Databaseはこの選択肢では単一ライターであり、欧州からの通常時ローカル書き込みを満たさない。
3. DynamoDB Global Tablesは各レプリカで読み書きできるマルチリージョン・マルチアクティブをマネージドに提供する。

4 正解が要件を満たす理由

Global Tablesでは両リージョンのレプリカがローカルの読み取りと書き込みを処理し、変更を他リージョンへ自動複製します。アプリケーションが複製遅延と競合を処理できるという条件があるため、マルチアクティブのトレードオフも受容できます。Route 53のレイテンシーおよびヘルスチェックと組み合わせることで、通常時の低遅延とリージョン障害時の継続性を、独自複製基盤なしで実現します。

5・6 各選択肢の検討

不適切な理由と、正解になり得る条件を対にして確認します。

A 不正解

PITRからの別リージョン復元は復旧時間を要し、通常時の欧州ローカル書き込みも提供しません。

この選択肢が正解になる条件

RTOが数時間で、平常時のDRリージョン費用を最小にし、過去時点への論理復旧が主目的なら適切です。

B 不正解

S3へのイベント蓄積とインポートはオンラインの双方向複製ではなく、障害時のテーブル作成が数分のRTOを危うくします。

この選択肢が正解になる条件

DynamoDBデータを分析用データレイクへ継続配信し、復旧は長時間でもよいオフラインアーカイブ要件なら候補です。

C 正解

各リージョンで読み書き可能なレプリカとマネージド複製が、active-active、低遅延、低運用負荷を同時に満たします。

D 不正解

提示構成は北米だけがライターであり、欧州からの通常時ローカル書き込みという核心要件を満たしません。データモデル変更も必要です。

この選択肢が正解になる条件

リレーショナルSQLが必須で、通常時は単一リージョン書き込み、他リージョンは低遅延読み取りとDR昇格でよいなら適切です。

7 今後使える判断基準

マルチリージョンDBでは、各リージョンが通常時に書けるか、複製整合性と競合を許容できるかを先に確認する。両方に書くNoSQLで競合許容ならGlobal Tables、単一ライターであれば他のグローバルDBも比較する。

8 関連サービスとの比較

DynamoDB Global Tablesはマルチアクティブの読み書き、PITRは過去時点への復元、Streamsは変更イベント処理です。Aurora Global DatabaseはリレーショナルなクロスリージョンDRと読み取りに強い一方、典型構成は単一プライマリライターです。

QUESTION 074

問題74

PERFORMANCE

本番より難しい

単一選択

CloudFrontキャッシュキーとOrigin Shield

旅行会社は世界向けの商品検索結果ページをCloudFront経由で配信し、オリジンは単一リージョンのALBです。URLには検索語、通貨、追跡用パラメータが含まれ、全リクエストにセッションCookieと多種類のUser-Agentが付きます。応答内容を変えるのは正規化済みの検索語と通貨だけで、追跡パラメータ、セッションCookie、User-Agentでは変わりません。人気検索の結果は5分間同じですが、現状はすべてのCookie、ヘッダー、クエリ文字列をキャッシュキーに含めているためヒット率が8%しかなく、多数のエッジから同時にオリジンへ同じ問い合わせが到達します。ピーク時には同一検索の要求が数千件重なります。個人情報誤って共有せず、鮮度5分を守りながら、世界規模のオリジン負荷と応答遅延を最も改善する構成はどれですか。

- A** すべてのCookie、ヘッダー、クエリ文字列をオリジンへ転送し、同じ値をキャッシュキーに維持する。TTLだけを24時間へ延長して全体のキャッシュヒット率を上げる。
- B** キャッシュを無効化し、CloudFront Functionsで全リクエストを単一URLへ書き換える。ALBのスティッキーセッションにより同じ利用者を同じEC2へ送る。
- C** 最小TTLを1時間に設定し、オリジンがno-cacheを返してもCloudFrontへ保存する。Cookieだけをキーから除外し、全ヘッダーと全クエリは維持する。
- D** キャッシュポリシーで正規化済み検索語と通貨だけをキャッシュキーに含め、不要なCookie・ヘッダー・追跡クエリを除外する。TTLを5分にし、ALBに近いリージョンでOrigin Shieldを有効化する。

ここで答えを決める

正解・解説は次のページから

ANSWER 074

問題74の正解・解説

1 正解 **D**

- D** キャッシュポリシーで正規化済み検索語と通貨だけをキャッシュキーに含め、不要なCookie・ヘッダー・追跡クエリを除外する。TTLを5分にし、ALBに近いリージョンでOrigin Shieldを有効化する。

2 問題文の重要な条件

- 応答を変える値は検索語と通貨だけで、個人別データは含まれない
- 許容鮮度は5分であり、それを超えるキャッシュは不可
- 多数のエッジから同一オリジンへの重複要求を集約したい

3 正解へ至る判断過程

1. キャッシュキーは応答表現を変える入力だけに限定しなければヒット率が細分化される。
2. TTLは鮮度要件の5分に合わせ、長すぎる固定最小TTLを避ける。
3. Origin Shieldを単一の追加キャッシュ層として置くと、多数のエッジからのオリジンミスをさらに集約できる。

4 正解が要件を満たす理由

Dは同じ応答を生成するリクエストを同一キャッシュキーへ集約し、個人情報に影響しない値だけを除外するため、安全にヒット率を高めます。5分TTLで鮮度を守り、Origin Shieldが各エッジやリージョナルキャッシュからの重複オリジンフェッチを集約します。これによりALBへの問い合わせ数と世界各地のミス時遅延を同時に減らせます。

5・6 各選択肢の検討

不適切な理由と、正解になり得る条件を対にして確認します。

A 不正解

キーの高カーディナリティを解消せず、24時間TTLは5分の鮮度要件にも反します。

この選択肢が正解になる条件

各Cookie・ヘッダー・クエリの組み合わせで本当に内容が変わり、結果が24時間不変であるなら正しい設計になり得ます。

B 不正解

キャッシュ無効ではオリジン負荷が減らず、異なる検索語や通貨を単一URLへまとめると正しい結果を返せません。

この選択肢が正解になる条件

全リクエストが同じ内容で、処理自体はキャッシュ不可だが接続のセッション維持だけが必要なステートフルアプリなら一部要素を検討できます。

C 不正解

1時間の最小TTLは5分の鮮度を破り、全ヘッダーと追跡クエリを残すためヒット率も低いままです。

この選択肢が正解になる条件

コンテンツが最低1時間不変で、全ヘッダー・クエリが表現を変えるという明確な要件がある場合に限り成立します。

D 正解

表現を変える値だけのキー、要件内TTL、追加キャッシュ層により正確性とオリジン保護を両立します。

7 今後使える判断基準

CloudFrontでは、キャッシュキーに含める値とオリジンへ転送する値を要件から分離する。表現を変えない高カーディナリティ値をキーから除き、TTLを鮮度に合わせ、多数エッジのミス集約にはOrigin Shieldを検討する。

8 関連サービスとの比較

キャッシュポリシーはキーとTTLを制御し、オリジンリクエストポリシーはキャッシュキーに入れずに転送する値を制御します。Origin Shieldは追加の集中キャッシュ層であり、Global Acceleratorのような非キャッシュ型の経路最適化とは異なります。

QUESTION 075

問題75

COST

複合難問

単一選択

請求分析・明細・通知の使い分け

200アカウントを持つ企業はOrganizationsの一括請求を使用しています。財務部は月次請求が予算を15%超過した原因を、リンクアカウント、サービス、使用タイプ、コスト配分タグ、Savings Plansの適用状況まで掘り下げ、SQLで独自の配賦レポートを長期保存したいと考えています。監査のため元データをS3に保持し、月ごとに再集計できることも必要です。各プロダクト責任者はコンソールで直近12か月の傾向と今後の予測を対話的に確認したいと考えています。また、月末の請求確定を待たず、実績または予測コストが部門別しきい値を超えそうなときに通知する必要があります。毎分のリアルタイム遮断や自動停止は不要で、数時間単位の更新で構いません。これら3種類の要件を、目的に合うマネージド機能で最も適切に満たす構成はどれですか。

- A** AWS Cost and Usage ReportをS3へ配信してAthenaで明細分析し、Cost Explorerで傾向・予測を確認する。コスト配分タグ等で絞ったAWS Budgetsを作成し、実績および予測しきい値の通知を設定する。
- B** AWS Budgetsだけを使用し、各予算の通知メールを月末にCSVとして結合する。通知履歴をSQL明細の代わりにし、将来予測にはCloudWatchメトリクスを使用する。
- C** Cost Explorerの画面だけを監査記録として長期保存し、スクリーンショットから使用タイプ別SQLテーブルを生成する。超過通知はCost Explorerの月次レポートメールだけで行う。
- D** CloudTrail organization trailをS3へ集約し、AthenaでAPI呼び出し回数を料金へ換算する。サービスクォータの使用率アラームを、部門予算超過の通知として利用する。

ここで答えを決める

正解・解説は次のページから

ANSWER 075

問題75の正解・解説

1 正解 **A**

- A** AWS Cost and Usage ReportをS3へ配信してAthenaで明細分析し、Cost Explorerで傾向・予測を確認する。コスト配分タグ等で絞ったAWS Budgetsを作成し、実績および予測しきい値の通知を設定する。

2 問題文の重要な条件

- SQLで扱う最詳細レベルの長期請求データと配賦が必要
- 責任者には対話的な傾向分析と予測が必要
- 実績または予測の部門別しきい値通知が必要だがリアルタイム制御ではない

3 正解へ至る判断過程

1. 詳細明細の機械処理と長期保存にはS3へ配信できるCost and Usage Reportを割り当てる。
2. グラフ、フィルター、過去傾向、予測の対話分析にはCost Explorerを割り当てる。
3. 金額・使用量・予約やSavings Plansなどのしきい値と通知にはAWS Budgetsを割り当てる。

4 正解が要件を満たす理由

Aは一つのツールに無理に集約せず、CURを詳細な請求データ基盤、Cost Explorerを対話分析、Budgetsをしきい値監視として使い分けます。CURはS3上の明細をAthenaでSQL処理してタグや割引適用を含む配賦に利用でき、Cost Explorerは傾向と予測を迅速に可視化し、Budgetsは実績・予測コストに基づく部門別通知を提供します。

5・6 各選択肢の検討

不適切な理由と、正解になり得る条件を対にして確認します。

A 正解

明細、対話分析、しきい値通知をそれぞれ最適なサービスへ割り当て、全条件を満たします。

B 不正解

Budgetsはしきい値管理用で、使用タイプや割引適用を含む行レベル明細のSQL分析基盤にはなりません。CloudWatchも請求予測の代替ではありません。

この選択肢が正解になる条件

必要なのが部門別の単純なしきい値通知だけで、明細分析や独自配賦が不要ならBudgets中心で十分です。

C 不正解

Cost Explorerは対話分析に適しますが、スクリーンショットは正規化された詳細請求データやSQL長期保存にならず、月次メールだけでは予測超過通知も不十分です。

この選択肢が正解になる条件

少数アカウントの概況を人が対話的に確認するだけで、明細エクスポートや自動しきい値通知が不要なら利用できます。

D 不正解

CloudTrailはAPI監査ログであり、実際の使用量、料金、割引、税、クレジットを表す請求明細ではありません。クォータ使用率も予算とは別物です。

この選択肢が正解になる条件

誰がどのAPIでリソースを変更したかを追跡し、サービス上限接近を監視することが目的ならこの2機能の組み合わせが適切です。

7 今後使える判断基準

コスト管理では、詳細データ取得、探索・可視化、しきい値通知を分ける。CURはデータセット、Cost Explorerは分析画面、Budgetsは計画値との比較と通知、と役割で判断する。

8 関連サービスとの比較

CURはS3に行レベルの請求・使用データを配信しAthena等で加工できます。Cost Explorerは管理された可視化と予測、AWS Budgetsは実績・予測のしきい値通知に向きます。CloudTrailやCloudWatchは監査・運用メトリクスで、請求明細の代替ではありません。

QUESTION 076

問題76

SECURE

本番相当

複数選択 (2つ)

S3 Access Pointsによる部門分離

大手小売企業は、分析用の単一S3バケットに営業、物流、財務のデータをプレフィックス別に保存しています。各部門は別AWSアカウントで処理を実行し、営業はsales/の読み書き、物流はlogistics/の読み書き、財務はfinance/の読み取りだけを必要とします。現在の巨大なバケットポリシーは変更のたびに競合し、誤って他部門のプレフィックスを許可する事故が起きました。セキュリティ部門はデータをコピーせず、部門ごとに独立して管理できるポリシーとエンドポイントを用意し、一部部門は指定VPCからだけアクセス可能にしたいと考えています。既存のオブジェクト所有権とライフサイクルは単一バケットで維持します。また、意図しないバケット名での直接アクセスにより部門別制御が迂回されないようにする必要があります。実施すべき対応を2つ選択してください。

- A** 部門ごとにS3 Access Pointを作成し、対象プレフィックス、許可操作、部門アカウントのプリンシパルを各アクセスポイントポリシーで制限する。必要な部門のアクセスポイントはVPCオリジンに限定する。
- B** 基盤バケットのポリシーで、承認済みアクセスポイントを経由する要求だけを許可するよう委任または制限し、部門ロールのIAMポリシーも対応するアクセスポイントARNに限定する。
- C** 部門ごとに新しいS3バケットを作成し、既存バケットから同じオブジェクトをCRRで継続複製する。各コピーのACLを部門アカウントへ付与し、元バケットは全社読み取り可能にする。
- D** 単一のS3 Access Pointを全社で共有し、そのポリシーでは全プレフィックスへのアクセスを許可する。部門分離は各オブジェクトACLだけで管理し、直接バケットアクセスも維持する。

ここで答えを決める

正解・解説は次のページから

ANSWER 076

問題76の正解・解説

1 正解 **A・B**

- A** 部門ごとにS3 Access Pointを作成し、対象プレフィックス、許可操作、部門アカウントのプリンシパルを各アクセスポイントポリシーで制限する。必要な部門のアクセスポイントはVPCオリジンに限定する。
- B** 基盤バケットのポリシーで、承認済みアクセスポイントを経由する要求だけを許可するよう委任または制限し、部門ロールのIAMポリシーも対応するアクセスポイントARNに限定する。

2 問題文の重要な条件

- 単一データセットをコピーせず部門ごとにポリシーを分離する
- プレフィックス、操作、アカウント、一部VPCという異なる境界がある
- バケットへの直接アクセスでAccess Point制御を迂回させない

3 正解へ至る判断過程

1. S3 Access Pointsは同じバケットに複数の名前付きエンドポイントと個別ポリシーを提供するため、部門単位の管理に合う。
2. アクセスポイントポリシーだけでなく、基盤バケットと呼び出し元IAMの双方でも要求を許可し、直接バケット経路を閉じる必要がある。
3. コピーとACLはデータ重複、整合性、分散した権限管理を増やし、目的に反する。

4 正解が要件を満たす理由

Aで部門ごとの名前付きエンドポイントに、プレフィックス、操作、プリンシパル、ネットワークオリジンを局所的に定義できます。Bで基盤バケットと部門ロール側もアクセスポイント経由に限定するため、直接バケットARNを使用した迂回を防げます。オブジェクトを複製せず、部門ポリシーを独立に変更でき、最小権限と運用分離を両立します。

5・6 各選択肢の検討

不適切な理由と、正解になり得る条件を対にして確認します。

A 正解

アクセスポイントごとの権限とネットワーク制御により、同一バケット上で部門別の入口を分離できます。

B 正解

Access Pointと基盤バケットのポリシーは組み合わせて評価されるため、バケット側の委任・制限とIAM側ARN制限が迂回防止になります。

C 不正解

3つのコピーの保存・複製費用、複製遅延、ACL管理を増やし、データをコピーしない条件に反します。元バケット全社読み取りも分離を破ります。

この選択肢が正解になる条件

部門ごとにリージョン、保持期間、所有権、暗号化キーを完全分離し、複製費用と整合性遅延を許容する要件なら別バケットが適切です。

D 不正解

単一の全許可アクセスポイントと直接アクセスでは権限変更の競合も迂回も解消せず、ACL中心管理は規模拡大に不向きです。

この選択肢が正解になる条件

全利用者が同じデータと操作権限を持ち、ネットワーク制限も共通で、部門分離が不要なら単一アクセスポイントで足ります。

7 今後使える判断基準

共有S3データセットの利用者ごとに権限・ネットワーク条件が違う場合はAccess Pointで入口を分割し、基盤バケットポリシーとIAMポリシーも含む三者の許可関係と直接アクセス経路を確認する。

8 関連サービスとの比較

S3 Access Pointsは同一バケットへの個別エンドポイントとリソースポリシーを提供します。別バケットは所有権・ライフサイクル・リージョンまで分けたい場合、Object ACLは個別オブジェクトの例外的管理であり、大規模な部門分離にはポリシー中心設計が適します。

QUESTION 077

問題77

RESILIENT

本番より難しい

単一選択

DRリージョンの起動前提準備

製造会社は、80台のEC2とALB、Auto Scaling、EBSを使用するシステムのパイロットライトDRを別リージョンに設計しました。データスナップショットは既に30分ごとにコピーされています。しかし初回訓練では、DRリージョンのOn-Demand vCPUクォータが不足し、希望インスタンスタイプの起動に失敗しました。さらに、起動テンプレートが参照するAMIはソースリージョンにしかなく、手作業で作ったセキュリティグループには本番との差分があり、RTO 60分を5時間超過しました。復旧先のサブネットとIAM設定にも更新漏れがありました。平常時に80台を常時稼働する予算はなく、予約済みの二重容量も避けます。次回以降、容量を常時二重化せずに、再現性のある復旧をRTO内で実行するための最も包括的な対策はどれですか。

- A** 障害発生後にService QuotasへvCPU引き上げを申請し、AMIをオンデマンドでコピーする。セキュリティグループは直近の本番画面を見ながら手動で作成する。
- B** DRリージョンで必要なvCPU等のクォータを事前確認・引き上げし、暗号化AMIとスナップショットを継続コピーする。全インフラをIaC化して定期デプロイ試験と復旧手順の自動テストを行う。
- C** スナップショットのコピー頻度を5分へ短縮するだけでRTOを改善する。コンピューティングとネットワークは障害発生後に運用担当者がコンソールからすべて手動作成する。
- D** DRリージョンに1台の踏み台EC2を起動し、そこへ本番AMIのIDとセキュリティグループ設定をテキストで保存する。必要な残り79台は障害時に同じ設定を参照して作る。

ここで答えを決める

正解・解説は次のページから

ANSWER 077

問題77の正解・解説

1 正解 **B**

- B** DRリージョンで必要なvCPU等のクォータを事前確認・引き上げし、暗号化AMIとスナップショットを継続コピーする。全インフラをIaC化して定期デプロイ試験と復旧手順の自動テストを行う。

2 問題文の重要な条件

- 失敗原因はデータだけでなく、クォータ、リージョン固有AMI、構成ドリフト
- RTO 60分で80台を再現性高く作る必要がある
- 平常時の全容量稼働は予算上不可

3 正解へ至る判断過程

1. RPOを改善するスナップショット頻度だけでは、起動不能と手作業によるRTO超過を解決しない。
2. クォータ引き上げとAMIコピーはリージョンごとに障害前に準備する必要がある。
3. IaCと定期試験によりネットワーク、セキュリティ、起動テンプレートを再現し、ドリフトと手作業時間を復旧経路から除く。

4 正解が要件を満たす理由

Bは訓練で判明した3つの前提条件をすべて事前に解消します。DRリージョンのクォータを確保し、リージョン固有のAMIと暗号化スナップショットを利用可能にし、IaCでALB、Auto Scaling、ネットワーク、セキュリティ設定を同じ定義から再構築します。定期的な実デプロイと復旧試験でクォータや参照先の陳腐化も検出でき、80台を常時稼働せずにRTO達成可能性を高めます。

5・6 各選択肢の検討

不適切な理由と、正解になり得る条件を対にして確認します。

A 不正解

クォータ承認とAMIコピーを障害後に始めると所要時間が不確定で、手動構築によるドリフトも再発します。

この選択肢が正解になる条件

RTOが数日で、復旧頻度が極端に低く、手作業と承認待ちを明示的に許容する非重要環境なら成立し得ます。

B 正解

容量上限、リージョン資産、構成再現性、継続検証を障害前に準備し、常時コンピュート費用を避けます。

C 不正解

コピー頻度はRPOを短縮しますが、クォータ不足、AMI欠如、手動ネットワーク構築というRTOの原因を解消しません。

この選択肢が正解になる条件

唯一の課題が許容データ損失で、インフラは既に稼働・検証済みであるならコピー頻度短縮が正解になります。

D 不正解

テキスト記録と踏み台は実行可能な構成定義でも容量保証でもなく、79台と周辺リソースを手動作成するためRTOを安定して満たせません。

この選択肢が正解になる条件

数台だけの小規模な検証環境で、手動復旧訓練が目的かつ長いRTOを許容するなら簡易ランブックとしては使えます。

7 今後使える判断基準

DR準備ではデータコピーだけでなく、復旧先リージョンのクォータ、リージョン固有アーティファクト、KMS権限、IaC、定期テストを一組で確認する。RTOは障害後に残した最も遅い未準備作業で決まる。

8 関連サービスとの比較

AWS Backupやスナップショットコピーはデータ復旧を担い、AMIコピーは起動イメージを用意します。Service Quotasは起動可能量、CloudFormation等のIaCは構成再現性を担うため、いずれか一つだけでは完全なDR準備になりません。

QUESTION 078

問題78

PERFORMANCE

複合難問

単一選択

大帯域ハイブリッド接続の冗長化

研究機関はオンプレミスの2拠点からAWSへ毎日40TBの測定データを送り、処理中は継続して2Gbps以上の安定スループットを必要とします。公共インターネット経由のSite-to-Site VPNだけで試験したところ、暗号化オーバーヘッドと経路変動により処理期限を超えました。主要拠点のルーター、通信事業者回線、またはDirect Connectロケーションのいずれか一つが失われても通信を継続し、計画保守時には自動的に別経路へ切り替える必要があります。切り替えのためのDNS変更や手動操作は許容されません。機密データのためトラフィック暗号化も必要です。通常時の性能には専用接続を使い、広域障害時には性能が低下してもインターネット経由の最終バックアップを残したいと考えています。最も要件に合う接続設計はどれですか。

- A** 1本の10Gbps Direct Connect専用接続を1ロケーションに用意し、単一のプライベートVIFを使用する。Direct Connect自体が冗長なため追加経路は作らない。
- B** 各拠点から1つずつSite-to-Site VPN接続を作成し、各接続の2トンネルをBGPで利用する。Direct Connectは使用せず、VPNを追加すれば常時2Gbpsを保証できると見なす。
- C** 異なるDirect Connectロケーションと顧客ルーターを使う冗長な専用接続を構成し、BGPでフェイルオーバーを設計する。必要に応じMACsecまたはVPNで暗号化し、別経路のSite-to-Site VPNも最終バックアップにする。
- D** 各オンプレミス拠点とAWS VPCをVPC peeringで直接接続し、Route 53ヘルスチェックで正常なピアへ切り替える。データ転送にはS3 Transfer Accelerationを併用する。

ここで答えを決める

正解・解説は次のページから

ANSWER 078

問題78の正解・解説

1 正解 C

- C** 異なるDirect Connectロケーションと顧客ルーターを使う冗長な専用接続を構成し、BGPでフェイルオーバーを設計する。必要に応じMACsecまたはVPNで暗号化し、別経路のSite-to-Site VPNも最終バックアップにする。

2 問題文の重要な条件

- 通常時に2Gbps以上の安定した大容量転送が必要
- 回線、ルーター、Direct Connectロケーションの単一障害を排除する
- 暗号化と、インターネット経由の最終バックアップを両立する

3 正解へ至る判断過程

1. VPNだけではインターネット経路の変動を残し、試験でも性能要件を満たしていない。
2. Direct Connectを1本だけにすると、ポート、ルーター、回線、ロケーションが単一障害点になる。
3. 異なるロケーション・装置のDirect ConnectをBGPで冗長化し、暗号化層と独立VPNを重ねることで性能、機密性、最終退避を分担できる。

4 正解が要件を満たす理由

Cは高帯域で予測可能な通常経路をDirect Connectにし、接続を異なるロケーションと顧客装置へ分散して単一障害点を排除します。BGPで経路優先度と自動切り替えを構成でき、接続方式と要件に応じてMACsecまたはDirect Connect上を含むVPNで暗号化できます。さらに独立したSite-to-Site VPNを残すため、Direct Connect側の広域障害でも性能を落として接続を継続できます。

5・6 各選択肢の検討

不適切な理由と、正解になり得る条件を対にして確認します。

A 不正解

帯域は満たしても1接続、1ロケーション、1ルーターが単一障害点で、暗号化と最終バックアップも示されていません。

この選択肢が正解になる条件

非重要な開発転送で停止を許容し、単一接続障害時は復旧を待てて、別途暗号化不要なら低コスト案として成立します。

B 不正解

VPNの2トンネルは可用性を高めますが、公共インターネット経路で安定2Gbpsを保証する設計ではなく、既に性能試験で不合格です。

この選択肢が正解になる条件

必要帯域が各トンネル能力内で、経路変動を許容し、短期間で低コストの暗号化接続を作ることが優先ならVPNが適切です。

C 正解

専用接続の性能、物理・場所の冗長性、暗号化、独立VPNバックアップをすべて明示的に満たします。

D 不正解

VPC peeringはVPC間接続であり、オンプレミス拠点を直接ピア接続するサービスではありません。Transfer Accelerationも汎用ハイブリッドネットワークの代替ではありません。

この選択肢が正解になる条件

接続対象が重複しないCIDRを持つ2つのVPCで、非推移の対一プライベート通信だけが必要ならVPC peeringが適切です。

7 今後使える判断基準

ハイブリッド接続は帯域だけでなく、物理回線、顧客装置、AWSロケーション、論理セッションの障害ドメインを列挙する。安定大帯域はDirect Connect、暗号化はMACsec/VPN、独立退避は別VPNなど役割を分ける。

8 関連サービスとの比較

Direct Connectは専用接続による安定帯域、Site-to-Site VPNはインターネット上の迅速で暗号化された接続です。両者は代替だけでなく主系と退避系として併用できます。VPC peeringはVPC同士、Transfer AccelerationはS3オブジェクト転送に用途が限定されます。

QUESTION 079

問題79

COST

本番相当

単一選択

変動する長時間コンテナジョブの実行基盤

映像配信会社は、取引先から届く動画を変換するDockerコンテナをオンプレミスで実行している。1ジョブは4vCPU、8GBメモリを使用し、平均35分、最長50分かかる。ジョブ数は通常1時間に0～5件だが、新作公開時には200件が一度に到着し、深夜は8時間以上ゼロになる。ジョブ間でローカル状態を共有せず、入力と出力はS3に置ける。処理開始まで2分は許容できるが、開始したジョブの中断は認められない。既存コンテナを分割する改修は次の四半期までできず、少人数の運用チームはOSのパッチ、インスタンスの容量計画、クラスターのスケールインを担当したくない。月間の総実行時間はピーク用容量を常設するほど多くなく、ライセンスもホスト固定ではない。常時稼働容量への支払いを避けつつ、現行イメージを最小変更で移行する最も費用対効果の高い構成はどれですか。

- A** コンテナイメージをLambdaへ登録し、各ジョブを1回の関数呼び出しで処理する。メモリを最大まで増やし、タイムアウトは変更せず高い同時実行数でバーストを吸収する。
- B** ECS on EC2クラスターにピーク時200ジョブ分のリザーブドインスタンスを常時配置し、最低台数は維持したままSQSのキュー長に応じてECSタスクだけを増減する。
- C** ECS on EC2をSpotインスタンスだけのAuto Scalingグループで構成し、キャパシティリバランシングにより中断されたジョブを途中から必ず再開する。
- D** SQSにジョブを格納し、キュー長に応じてECSのオンデマンドFargateタスクを起動する。各タスクは1ジョブ完了後に終了し、一時的な失敗時はキューから安全に再試行する。

ここで答えを決める

正解・解説は次のページから

ANSWER 079

問題79の正解・解説

1 正解 **D**

- D** SQSにジョブを格納し、キュー長に応じてECSのオンデマンドFargateタスクを起動する。各タスクは1ジョブ完了後に終了し、一時的な失敗時はキューから安全に再試行する。

2 問題文の重要な条件

- 1ジョブが最長50分であり、既存コンテナを15分以内の処理へ分割できない
- 負荷はゼロから急増まで変動し、長いアイドル時間への常時課金を避けたい
- 開始済みジョブは中断不可で、サーバー管理と容量計画も最小化する

3 正解へ至る判断過程

1. Lambdaは従量課金と急な並列化には合うが、1回の実行上限を超えるため、分割不可という移行制約に反する。
2. EC2は定常的で高使用率なら単価を下げやすい一方、今回のピーク容量常設は大きなアイドル費用を生み、クラスター運用も残る。
3. 中断不能なのでSpotを主系にせず、ジョブの間だけ課金され、サーバー管理不要で長時間コンテナをそのまま動かせるオンデマンドFargateを選ぶ。

4 正解が要件を満たす理由

Dは既存の長時間DockerイメージをECSタスクとしてほぼそのまま実行でき、Lambdaの15分制限を受けません。SQSを作業の耐久性のある入口にして必要数だけFargateタスクを起動し、完了後に停止すれば、8時間以上の無負荷時間にコンピューティング料金を払い続けずに済みます。FargateはEC2クラスターのプロビジョニング、パッチ、配置最適化を不要にし、オンデマンド容量を使うためSpot中断の条件にも抵触しません。

5・6 各選択肢の検討

不適切な理由と、正解になり得る条件を対にして確認します。

A 不正解

Lambdaは1回の実行時間が最大15分であり、35～50分のジョブを分割せず完了できません。メモリや同時実行数を増やしても実行時間上限は変わりません。

この選択肢が正解になる条件

各動画を15分未満で処理できるか、Step Functionsなどで安全に小さな冪等ステップへ分割でき、呼出しが短時間かつ散発的ならLambdaが有力です。

B 不正解

ピーク容量を常設すると無負荷時間にもEC2料金が発生し、ホストOS、クラスター容量、スケールインの運用も残るため、負荷形状と運用要件に合いません。

この選択肢が正解になる条件

年間を通じて高いベース負荷がありインスタンスを高使用率で埋められ、Savings PlansやRIの割引と専用ホスト機能が運用費を上回るならECS on EC2が適します。

C 不正解

Spotは安価でも回収によりタスクが中断され得ます。キャパシティリバランシングは新容量の確保を助けますが、アプリケーションの途中状態を自動的かつ必ず復元する機能ではありません。

この選択肢が正解になる条件

ジョブがチェックポイントから再開でき、再試行による期限超過を許容し、割込み通知を処理できるならFargate SpotまたはEC2 Spotが最安候補になり得ます。

D 正解

長時間の既存コンテナ、中断不可、ゼロを含む変動負荷、サーバー運用回避を同時に満たし、実行しているタスク分だけ支払えます。

7 今後使える判断基準

実行基盤は平均時間だけでなく、最大実行時間、負荷のアイドル率、中断許容性、既存パッケージ、運用人員で選ぶ。15分以内の短いイベントはLambda、変動する長時間コンテナはFargate、定常高稼働や特殊ホスト要件はEC2が基本線となる。

8 関連サービスとの比較

Lambdaはリクエスト単位の短時間処理と高速な自動スケールに強い一方、実行時間上限があります。Fargateはサーバーなしで長時間コンテナをタスク単位課金でき、EC2はホスト制御と高い定常利用率での割引余地がある代わりに容量・OS運用を負います。Spotは中断許容ワークロード向けです。

QUESTION 080

問題80

SECURE

本番より難しい

複数選択 (2つ)

公開WebサービスのL7攻撃とDDoS対策

世界向けECサイトは、Route 53、CloudFront、インターネット向けALB、ECSで構成され、通常は毎秒8,000リクエスト、セール時は毎秒80,000リクエストを処理する。直近の攻撃では、SQLインジェクションを含むリクエストと、少数IPからログインAPIへの大量試行が混在し、別の日にはL3/L4のDDoSでALBのスケール費用が急増した。通常のセールの遮断せず、既知のWeb脆弱性とエンドポイント単位の過剰リクエストをマネージドに抑止したい。また経営陣は、大規模DDoS時の専門家による24時間支援、攻撃に伴う保護対象リソースのスケールアップ費用への追加保護、ヘルスに基づく迅速な検出を求めている。ルール導入時の誤検知を事前評価でき、アプライアンスの運用は避けることも条件である。要件を満たす対策を2つ選択してください。

- A** CloudFrontにAWS WAFのWeb ACLを関連付け、AWS Managed RulesをまずCountで評価してからBlockへ移す。ログインURIにscope-downしたrate-based ruleも追加する。
- B** ECSのVPCへAWS Network Firewallを配置し、すべての送信通信とALBからタスクへの通信を検査する。SQL文字列と送信元IP頻度はステートフルルールで遮断する。
- C** CloudFront、ALB、Route 53をShield Advancedの保護対象にし、Route 53ヘルスチェックと通知を構成する。事前にShield Response Teamへのアクセスも承認する。
- D** 自動適用されるShield Standardだけを使用し、GuardDutyの検出結果を受けて担当者がセキュリティグループへ攻撃元IPを随時追加して個別に対応する。

ここで答えを決める

正解・解説は次のページから

ANSWER 080

問題80の正解・解説

1 正解 A・C

- A** CloudFrontにAWS WAFのWeb ACLを関連付け、AWS Managed RulesをまずCountで評価してからBlockへ移す。ログインURIにscope-downしたrate-based ruleも追加する。
- C** CloudFront、ALB、Route 53をShield Advancedの保護対象にし、Route 53ヘルスチェックと通知を構成する。事前にShield Response Teamへのアクセスも承認する。

2 問題文の重要な条件

- SQLインジェクションと特定APIへの高頻度アクセスというL7制御が必要
- 正常な大規模セールを維持しながらルールを安全に導入する
- 大規模DDoS時の専門支援、ヘルス連携、費用保護まで求められる

3 正解に至る判断過程

1. HTTPリクエスト内容とURI別レートを評価する要件には、CloudFrontの手前でAWS WAFを適用し、Managed Rulesとrate-based ruleを組み合わせる。
2. WAFルールの初期導入はCountとログで誤検知を確認してからBlockへ移し、通常セールの大量トラフィックを単純な全体上限で遮断しない。
3. ネットワーク層のDDoS保護だけでなくSRT支援やDDoSコスト保護を明示しているため、Standardではなく対象リソースへShield Advancedを追加する。

4 正解が要件を満たす理由

AはCloudFrontで悪意あるHTTPリクエストをオリジン到達前に評価し、Managed Rulesで一般的なWeb攻撃を、scope-downしたrate-based ruleでログインAPIだけへの過剰試行を抑止します。Countから始めることで正常なセール通信への影響も検証できます。CはCloudFront、ALB、Route 53に高度なDDoS検出・緩和を追加し、ヘルスチェック連携、Shield Response Teamの支援、適格なDDoS関連スケーリング費用の保護というWAF単独では満たせない条件を補完します。

5・6 各選択肢の検討

不適切な理由と、正解になり得る条件を対にして確認します。

A 正解

L7の既知攻撃とURI単位のレート制御をエッジで実施し、Count評価によって正常トラフィックを守りながら導入できます。

B 不正解

Network FirewallはVPC境界のネットワーク検査には有効ですが、CloudFrontで受けるWebリクエストのWAF代替ではなく、オリジン到達前のURI別レート制御やShieldの費用保護も提供しません。

この選択肢が正解になる条件

主目的がVPC間・オンプレミス間・インターネット送信の集中検査、ドメインやシグネチャに基づくネットワーク制御で、HTTPエッジ保護は別途WAFが担うなら適切です。

C 正解

大規模DDoSに対する高度な保護、ヘルスベース検出、SRT支援、適格なDDoSコスト保護という経営要件を満たします。

D 不正解

Shield Standardは基礎的なDDoS保護を自動提供しますが、設問が求めるShield Advanced固有のSRT支援やコスト保護を満たしません。手動IP追加は分散・変動する攻撃に追従しにくく、L7内容も評価できません。

この選択肢が正解になる条件

小規模でリスクの低い公開サービスが基礎的なL3/L4保護だけを求め、専門支援やDDoSコスト保護を必要としないならShield Standardを前提にできます。

7 今後使える判断基準

Web攻撃対策では層を分ける。HTTP内容、パス、IP別頻度はAWS WAF、一般的なL3/L4 DDoSの基礎保護はShield Standard、高度な検出・SRT・コスト保護が必要ならShield Advancedを選び、WAFルールはCountで誤検知を較正する。

8 関連サービスとの比較

AWS WAFはSQLi、XSS、Bot、レート制御などL7のリクエスト選別を担当します。Shield Standardは自動の基本DDoS保護、Shield Advancedは保護対象への高度なDDoS機能と支援を追加します。Network FirewallはVPCのL3～L7ネットワークフロー検査であり、CloudFront Web ACLの置換ではありません。

QUESTION 081

問題81

RESILIENT

複合難問

単一選択

サーバーレス接続急増時のデータベース保護

保険会社はAPI GatewayとLambdaから、RDS for MySQLのMulti-AZ DBインスタンスへ短いトランザクションを実行している。平常時のLambda同時実行数は500だが、災害受付時には数分で8,000へ増える。DBのCPU使用率は40%以下でも接続数上限に達し、TLS接続と認証の繰り返しで応答が悪化する。フェイルオーバー訓練では、多数の関数が失効した接続を同時に再作成して再試行の波が生じ、復旧が遅れた。DBを接続数だけのために大型化せず、アプリケーションのSQLとドライバーをほぼ変更せずに、接続を再利用し、フェイルオーバーの影響を軽減したい。認証情報はSecrets Managerで管理し、プロキシサーバーのパッチ運用も避ける必要がある。書込み処理を読取り専用系へ逃がすことはできない。最適な改善策はどれですか。

- A** RDS Proxyを複数AZのサブネットで作成し、Secrets ManagerまたはIAM認証を構成する。Lambdaの接続先をProxyエンドポイントへ変更し、プール上限と借用タイムアウトを調整する。
- B** DBインスタンスを最大クラスへ変更してmax_connectionsを引き上げ、各Lambda呼び出しがDBエンドポイントへ新規接続する。失敗時は待ち時間なしで再試行する。
- C** 同一リージョンにRDSリードレプリカを追加し、書込みを含む全Lambda接続をレプリカへ振り分ける。障害時はレプリカが自動的に同期スタンバイとして切り替わると見なす。
- D** 2つのAZのEC2 Auto ScalingグループにProxySQLを構築し、NLBの背後へ置く。独自スクリプトで接続プール、ヘルスチェック、DBフェイルオーバー追従を管理する。

ここで答えを決める

正解・解説は次のページから

ANSWER 081

問題81の正解・解説

1 正解 **A**

- A** RDS Proxyを複数AZのサブネットで作成し、Secrets ManagerまたはIAM認証を構成する。Lambdaの接続先をProxyエンドポイントへ変更し、プール上限と借用タイムアウトを調整する。

2 問題文の重要な条件

- Lambdaの急増による接続ストームがCPU余力のあるDBを枯渇させている
- Multi-AZフェイルオーバー時の再接続集中も軽減する必要がある
- SQLの大幅変更、DBの過大化、自己管理プロキシの運用を避けたい

3 正解へ至る判断過程

1. 問題はクエリ処理能力より、短命なクライアント接続数と接続確立コストなので、まず接続多重化・プーリングを選ぶ。
2. 読取り専用レプリカは書き込み接続を受けられず、同期スタンバイの代替でもないため、接続ストームの直接解決にならない。
3. RDS Proxyはマネージドな高可用プロキシとして接続を共有し、DB切替を追跡しつつアプリケーション側接続を維持しやすいため、変更と運用負荷が最小になる。

4 正解が要件を満たす理由

Aでは数千のLambdaクライアント接続を、RDS Proxyが管理するより少数のDB接続へ多重化・再利用できるため、接続上限、TLS・認証コスト、再試行の集中からDBを保護できます。Proxyはフェイルオーバー先へ自動的に接続し直し、アプリケーション接続を維持しやすくするので復旧影響も短縮します。Secrets ManagerとIAM認証へ統合でき、EC2プロキシの構築、パッチ、AZ冗長化をチームが運用する必要もありません。

5・6 各選択肢の検討

不適切な理由と、正解になり得る条件を対にして確認します。

A 正解

マネージドな接続プールでLambdaの接続数を吸収し、認証管理とDBフェイルオーバー追従を小さな接続先変更で実現します。

B 不正解

大型化は接続上限を増やしても、接続生成コストと無制御な即時再試行を残します。CPUに余力がある今回、平常時の費用も増やす対症療法です。

この選択肢が正解になる条件

CPU・メモリ・I/O自体が継続的に飽和し、接続プール後も実クエリ処理能力が不足するなら、適切なDBクラスへのスケールアップが必要です。

C 不正解

リードレプリカは非同期の読取り拡張先で、書込みを処理できません。Multi-AZの同期スタンバイとは役割が異なり、接続再利用も提供しません。

この選択肢が正解になる条件

負荷の大半が多少のレプリケーション遅延を許容する読取りで、接続数ではなく読取りCPUやI/Oがボトルネックならリードレプリカへの分散が適切です。

D 不正解

接続プール自体は有効ですが、EC2、ProxySQL、NLB、パッチ、監視、フェイルオーバー追従を自社運用するため、運用最小化の条件でRDS Proxyに劣ります。

この選択肢が正解になる条件

RDS Proxy非対応のDBエンジン・プロトコルや特殊なSQLルーティングが必須で、専任チームがプロキシを高可用運用できるなら自己管理方式が候補になります。

7 今後使える判断基準

サーバーレスからRDBへの障害では、CPU不足と接続不足を分けて考える。短命接続の急増、認証オーバーヘッド、フェイルオーバー再接続が症状ならRDS Proxy、読取り処理量ならリードレプリカ、実計算能力ならDBサイズを検討する。

8 関連サービスとの比較

RDS Proxyは接続プール、多重化、Secrets Manager/IAM統合、フェイルオーバー追従をマネージドで提供します。リードレプリカは非同期の読取りスケール、Multi-AZは高可用性、大型DBクラスは単体処理能力の拡張です。自己管理プロキシは柔軟ですが運用責任が増えます。

QUESTION 082

問題82

PERFORMANCE

本番相当

単一選択

多数VPCのハブ接続と経路分離

ソフトウェア企業は1リージョン、8アカウントに45個のVPCを持ち、今後1年で150個まで増える見込みである。CIDRは重複していない。各VPCは共有サービスVPCのDNS、監視、ライセンスサーバーと通信し、Direct Connectで接続されたオンプレミスにも到達する必要がある。一方、開発VPC同士と本番VPC同士の横方向通信は原則拒否し、共有サービスとオンプレミスへの経路だけを環境別に許可したい。現在のVPC peeringフルメッシュは接続数と各VPCルート更新が急増し、新規VPC追加に数日かかっている。アプリケーションはIP通信を使用するため、個別サービス単位の公開へ変更できない。接続は中央ネットワークアカウントから統制し、参加アカウントの所有権は維持したい。高いスループットを維持しつつ、接続と経路の運用を最も簡素化する設計はどれですか。

- A** VPC peeringを継続し、全VPC間に一対一の接続を作る。各ルートテーブルへ全相手CIDRを登録し、オンプレミス経路は共有サービスVPCを経由して推移的に利用する。
- B** Transit Gatewayをネットワークアカウントに作成してAWS RAMで共有する。VPCとDirect Connect gatewayを接続し、環境別のTGWルートテーブルで伝播と到達先を分離する。
- C** 共有サービスVPCの全サーバーをNLBの背後に置き、各VPCへPrivateLinkのインターフェイスエンドポイントを作る。オンプレミスとVPC間の任意IP通信も同じエンドポイントで中継する。
- D** 共有サービスVPCに複数のVPNアプライアンスを配置し、各VPCからSite-to-Site VPNを終端する。BGPフルメッシュによりVPC間とDirect Connectへの経路を交換する。

ここで答えを決める

正解・解説は次のページから

ANSWER 082

問題82の正解・解説

1 正解 **B**

- B** Transit Gatewayをネットワークアカウントに作成してAWS RAMで共有する。VPCとDirect Connect gatewayを接続し、環境別のTGWルートテーブルで伝播と到達先を分離する。

2 問題文の重要な条件

- VPC数が45から150へ増え、一対一接続と個別ルート更新が限界
- 共有サービスおよびオンプレミスへの推移的なIP到達性が必要
- 環境間の横方向通信を経路レベルで分離したい

3 正解へ至る判断過程

1. VPC peeringは一対一かつ非推移なので、フルメッシュと共有VPC経由のオンプレミス中継を大規模に管理する要件へ適合しない。
2. PrivateLinkは特定サービスの一方向公開には優れるが、任意IP接続やオンプレミス経路のハブとしては用途が異なる。
3. Transit Gatewayをハブにし、複数のTGWルートテーブルへの関連付けと伝播を制御すれば、アタッチメント数に比例する接続と環境別分離を両立できる。

4 正解が要件を満たす理由

Bは各VPCを1つのTransit Gatewayへアタッチするハブアンドスポークへ変え、VPC数の二乗で増えるpeering接続を回避します。Direct Connect gatewayとの関連付けによりオンプレミス経路も同じハブへ統合できます。開発、本番、共有サービス用のTGWルートテーブルを分け、アタッチメントの関連付けと経路伝播を選べば、スポーク同士を遮断しながら必要な共通宛先だけを公開できます。RAM共有によりアカウント分離も維持できます。

5・6 各選択肢の検討

不適切な理由と、正解になり得る条件を対にして確認します。

A 不正解

150 VPCのフルメッシュは接続とルートが急増し、VPC peeringは推移的ルーティングをサポートしないため、共有VPC経由で別VPCやオンプレミスへ中継できません。

この選択肢が正解になる条件

接続するVPCが2〜3個に限定され、低遅延・高帯域の一对一通信だけで必要で、推移的接続や集中ルート管理が不要ならVPC peeringが簡潔です。

B 正解

多数VPCとDirect Connectを単一ハブへ集約し、TGWルートテーブルで到達性をセグメント化できるため、拡張性と運用性を満たします。

C 不正解

PrivateLinkはNLB等の背後の特定サービスをプライベート公開する仕組みで、全CIDRへの双方向通信やオンプレミス経路の汎用トランジットを提供しません。

この選択肢が正解になる条件

利用側VPCから少数のプロバイダーサービスだけへ一方向に到達させ、CIDR重複を許容しつつネットワーク全体を接続したくないならPrivateLinkが適します。

D 不正解

実現可能な部分はあっても、VPN装置、トンネル、BGP、可用性、パッチをVPCごとに運用し、不要な暗号化オーバーヘッドと多数接続を持ち込むためマネージドTGWより複雑です。

この選択肢が正解になる条件

特殊なルーティング・暗号機能を持つ既存SD-WAN製品が必須で、Transit Gateway Connect等との統合を専任ネットワークチームが運用するならアプライアンスが必要になる場合があります。

7 今後使える判断基準

VPC接続は規模と通信モデルで選ぶ。少数の一对一・非推移通信はpeering、多数VPCとオンプレミスの推移的ハブはTransit Gateway、特定サービスだけを消費側へ公開するならPrivateLinkを使う。

8 関連サービスとの比較

VPC peeringは直接の一对一接続で処理ハブを持ちませんが非推移です。Transit Gatewayはアタッチメント課金とデータ処理料金の代わりに、ハブ型接続、複数ルートテーブル、VPN・Direct Connect統合を提供します。PrivateLinkはサービス単位の分離とCIDR重複への強さが特徴です。

QUESTION 083

問題83

COST

本番より難しい

単一選択

EFSの低頻度データとAZ耐障害性

設計事務所は3つのAZで稼働するECSサービスから、18TBの設計ファイルをNFSで共有している。調査では、作成後14日以内のファイルが頻繁に使われる一方、全容量の92%はその後ほとんど参照されず、月に数回だけ過去案件が開かれる。アプリケーションはPOSIXファイルロックと既存パスに依存するため、S3 APIへの変更はできない。現在のEFS Standardのストレージ費が高いが、AZ全体が停止しても別AZのタスクから同じファイルシステムへ継続アクセスできることが契約要件であり、バックアップから数時間かけて復元する方式は通常障害時の代替にならない。まれに再アクセスしたファイルには継続的に高速アクセスする。読取り時のデータアクセス料金は許容できる。可用性を下げず、運用を増やさずにストレージ費を削減する最適な構成はどれですか。

- A** EFS One Zoneを単一AZに作成し、14日後にOne Zone-IAへ移行する。別AZからはクロスAZでマウントし、AZ障害時はAWS Backupから新規EFSへ復元する。
- B** EFS Regionalを維持し、全ファイルをStandardストレージクラスに固定する。Provisioned Throughputを最大にして、低頻度ファイルの読取り遅延をなくす。
- C** EFS Regionalを維持し、最終アクセスから14日後にIAへ移すライフサイクルポリシーを設定する。再アクセス時にStandardへ戻すポリシーも有効にする。
- D** 全データをS3 Standard-IAへ移し、各ECSタスクでS3バケットをNFSとしてマウントする。既存のPOSIXロックとファイル更新はS3が透過的に再現すると見なす。

ここで答えを決める

正解・解説は次のページから

ANSWER 083

問題83の正解・解説

1 正解 **C**

- C** EFS Regionalを維持し、最終アクセスから14日後にIAへ移すライフサイクルポリシーを設定する。再アクセス時にStandardへ戻すポリシーも有効にする。

2 問題文の重要な条件

- 92%が低頻度だが、NFS、POSIXパス、ファイルロックを変更できない
- AZ障害時にも復元待ちなしで別AZから継続アクセスする必要がある
- 低頻度化は最終アクセス後に判断でき、再び活性化したファイルは高速層へ戻したい

3 正解へ至る判断過程

1. プロトコル制約からEFSを維持し、S3への置換によるアプリケーション意味論の変更を避ける。
2. One Zoneは安価でも単一AZにデータを置くため、AZ喪失時に同一ファイルシステムを継続利用する契約要件を満たさない。
3. EFS RegionalのMulti-AZ耐障害性を保ち、ライフサイクル管理で未アクセスファイルだけをIAへ移し、再アクセス時にStandardへ戻す。

4 正解が要件を満たす理由

CはRegional EFSの複数AZにまたがる可用性を維持するため、1 AZの障害時にも別AZのマウントターゲットから同じファイルシステムへアクセスできます。EFSライフサイクル管理は内部の最終アクセス時刻を基に低頻度ファイルをIAへ自動移行し、再アクセス後にStandardへ戻す設定も可能です。NFS名前空間やPOSIX操作を変えず、容量の92%だけを低単価層へ置くので、バックアップ復元手順を通常フェイルオーバーへ持ち込まずに費用を下げられます。

5・6 各選択肢の検討

不適切な理由と、正解になり得る条件を対にして確認します。

A 不正解

One Zoneは単一AZ内にデータを保存します。別AZからマウントできても、保存先AZそのものの喪失に耐えるRegional構成ではなく、復元時間が契約条件に反します。

この選択肢が正解になる条件

データが再生成可能、またはAZ障害時にバックアップ復元のRTOを許容でき、Multi-AZ耐障害性よりストレージ単価を優先するならOne ZoneとOne Zone-IAが適します。

B 不正解

可用性は満たしますが、全18TBをStandardに固定して低頻度データの費用を削減しません。スループット増強はアクセス頻度に基づくストレージ料金問題の解決ではありません。

この選択肢が正解になる条件

ほぼ全ファイルが常時アクセスされ、IAのデータアクセス料金が節約額を上回るか、一貫して非常に高いスループットが主課題ならStandard中心が適切です。

C 正解

RegionalのAZ耐障害性とNFS互換性を維持しつつ、アクセス履歴に応じたIA移行・再昇格をマネージドに実施します。

D 不正解

S3はオブジェクトストレージであり、EFSと同じPOSIXファイルロックや更新意味論を透過的には提供しません。アプリケーション変更不可の条件に反します。

この選択肢が正解になる条件

アプリケーションをS3 APIとオブジェクト更新モデルへ変更でき、共有POSIXファイルシステムが不要で、大容量オブジェクトの低コスト保管が目的ならS3が有力です。

7 今後使える判断基準

EFS費用を下げるときは、まずファイル共有の意味論を維持するか、次にAZ喪失時のRTOを確認する。Multi-AZ継続が必要ならRegional+IAライフサイクル、単一AZ喪失を許容できる再生成可能データならOne Zoneを検討する。

8 関連サービスとの比較

EFS Regionalは複数AZにデータを保存し、Standard・IA・Archive間のライフサイクルを利用できます。EFS One Zoneは単一AZで低コストですがAZ喪失リスクを受け入れます。S3はさらに低コストなオブジェクト保存に向く一方、NFS/POSIX共有の直接代替ではありません。

QUESTION 084

問題84

SECURE

複合難問

複数選択 (2つ)

共有EFSのテナント分離と鍵管理

医療SaaS企業は、3つのECS Fargateサービスで1つのEFS Regionalファイルシステムを共有して費用を抑えたい。サービスごとに`/tenant-a`、`/tenant-b`、`/tenant-c`だけを見せ、コンテナイメージ内のUID/GID設定が誤っていても、指定したPOSIX IDとしてアクセスさせる必要がある。あるサービスのタスクロールやコンテナが侵害されても、別サービスのディレクトリをマウントできてはならない。また患者データは保存時と通信時に暗号化し、保存時の鍵は中央セキュリティチームが管理する同一アカウントのカスタマーマネージドKMSキーで失効・監査できることが監査要件である。サービスごとに別のEFSを作ることは避け、マネージド機能で最小権限を実装したい。実施すべき対策を2つ選択してください。

- A** サービスごとにEFS Access Pointを作成し、専用ルートディレクトリと強制POSIX UID/GIDを設定する。ファイルシステムポリシーとタスクIAMで対応Access Point経由のClientMountだけを許可する。
- B** サービスごとに異なるセキュリティグループを割り当て、同じEFSマウントターゲットのTCP 2049を全グループへ許可する。ディレクトリ分離はネットワーク層だけに委ねる。
- C** 各タスクの起動時スクリプトで共有EFSルートをマウントし、chmodとchownで担当ディレクトリを設定する。コンテナのrootユーザーには全ディレクトリの修復権限を残す。
- D** EFS作成時に中央セキュリティチーム管理の同一アカウントのカスタマーマネージドKMSキーで保存時暗号化を有効にする。マウントヘルパーのTLSとIAM認証を使い、各ポリシーも限定する。

ここで答えを決める

正解・解説は次のページから

ANSWER 084

問題84の正解・解説

1 正解 **A・D**

- A** サービスごとにEFS Access Pointを作成し、専用ルートディレクトリと強制POSIX UID/GIDを設定する。ファイルシステムポリシーとタスクIAMで対応Access Point経由のClientMountだけを許可する。
- D** EFS作成時に中央セキュリティチーム管理の同一アカウントのカスタマーマネージドKMSキーで保存時暗号化を有効にする。マウントヘルパーのTLSとIAM認証を使い、各ポリシーも限定する。

2 問題文の重要な条件

- 1つのEFS内でサービスごとにルートディレクトリとPOSIX IDを強制する
- 侵害されたタスクが別Access Pointやファイルシステムルートを使う経路もIAMで閉じる
- 保存時は指定CMK、通信時はTLSという二つの暗号化要件がある

3 正解へ至る判断過程

1. セキュリティグループはNFSポートへの到達性を制御するだけで、同じファイルシステム内のディレクトリやPOSIX主体を分離しない。
2. EFS Access Pointなら、マウント時の見かけのルートと全NFS要求に用いるPOSIX IDをアプリケーション単位に強制できる。IAMとファイルシステムポリシーでAccess Point ARNも限定する。
3. 保存時暗号化はファイルシステム作成時にCMKを指定し、通信時暗号化はTLSマウントを使う。両者は別の制御なので一方で代用しない。

4 正解が要件を満たす理由

Aは各サービスに専用のアプリケーション入口を与え、Access Pointのルートディレクトリより上を見せず、クライアントが申告するUID/GIDではなく設定済みPOSIX IDでアクセスを評価させます。IAM認証とファイルシステムポリシーでタスクロールを対応Access Pointへ限定すれば、ルート直接マウントの迂回も抑止できます。Dは指定したカスタマーマネージドKMSキーでファイルデータを保存時暗号化し、メタデータはEFS内部管理キーで暗号化されます。TLSマウントで転送中も保護するため、単一EFSの費用効率を保ったまま分離と暗号化を完成させます。

5・6 各選択肢の検討

不適切な理由と、正解になり得る条件を対にして確認します。

A 正解

Access Pointが専用ルートとPOSIX IDを強制し、IAM・ファイルシステムポリシーとの組合せでサービス別のマウント経路を制限できます。

B 不正解

セキュリティグループはマウントターゲットへのTCP到達性を制御しますが、許可後にどのディレクトリを見られるか、どのUID/GIDとして操作するかは制御しません。

この選択肢が正解になる条件

ファイルシステム単位でネットワーク到達性を分けるだけでよく、各クライアントを同じ信頼境界として扱えるならセキュリティグループ制御が基本対策になります。

C 不正解

クライアント側スクリプトとroot権限を信頼する方式は、設定ミスやコンテナ侵害時に共有ルートへ到達でき、強制POSIX IDと迂回防止の要件を満たしません。

この選択肢が正解になる条件

全クライアントが同一の強い信頼境界にあり、単なる初期ディレクトリ作成を自動化するだけで、悪意あるクライアントからの分離が不要なら補助策として使えます。

D 正解

指定CMKによる保存時暗号化、TLSによる転送中暗号化、IAM認証を組み合わせ、監査可能な鍵統制を満たします。

7 今後使える判断基準

共有EFSの分離は、ネットワーク到達性、IAMによるマウント許可、Access Pointのルート/POSIX強制、暗号化を別レイヤーとして設計する。セキュリティグループだけでディレクトリ分離を実現したと判断しない。

8 関連サービスとの比較

EFS Access Pointは1ファイルシステム内のアプリケーション別入口とPOSIX強制を提供し、セキュリティグループはNFSネットワーク到達性を制御します。KMSは保存時の鍵、EFSマウントヘルパーのTLSは転送中暗号化、IAMとファイルシステムポリシーはClientMount等の認可を担当します。

QUESTION 085

問題85

RESILIENT

本番相当

単一選択

S3イベントの重複・順序逆転への耐性

不動産会社は、バージョニングを有効にしたS3バケットへ物件JSONをアップロードし、イベント通知でLambdaを起動してDynamoDBの検索用レコードを更新している。同じオブジェクトキーは担当者の修正で数秒間に複数回更新される。運用開始後、同じ作成イベントがまれに2回処理され、さらに新しい版の処理後に古い版の通知が届いてDynamoDBを過去の内容へ戻す事象が起きた。処理は通常数秒以内に開始し、Lambdaや下流の一時障害時には失わず再試行したい。全キーを1本の直列キューで処理してスループットを落とすことは避ける。更新処理はbucket、object key、version IDを取得でき、DynamoDBの条件付き書込みを追加できる。S3イベントの配送特性を前提に、重複と同一キー内の順序逆転を安全に処理する最適な構成はどれですか。

- A** S3からLambdaを直接呼び出し、関数の予約済み同時実行数を1にする。Lambdaが一度成功すればS3は同じイベントを再送せず、到着順もアップロード順になると見なす。
- B** S3イベントをEventBridgeから単一のSQS FIFOキューへ送り、すべてのオブジェクトに同じMessageGroupIdを設定する。FIFOの5分重複排除だけで永久に重複を防ぐ。
- C** S3イベントをEventBridgeアーカイブへ保存し、毎晩すべて再生する。DynamoDB更新は無条件PutItemとし、最後に再生されたイベントを最新版として扱う。
- D** S3通知をSQS StandardでバッファしてLambdaで処理する。bucket/keyごとにDynamoDBへ最大sequencerと処理IDを保持し、条件付き書込みで重複または古いイベントを拒否し、失敗は再試行・DLQへ送る。

ここで答えを決める

正解・解説は次のページから

ANSWER 085

問題85の正解・解説

1 正解 **D**

- D** S3通知をSQS StandardでバッファしてLambdaで処理する。bucket/keyごとにDynamoDBへ最大sequencerと処理IDを保持し、条件付き書込みで重複または古いイベントを拒否し、失敗は再試行・DLQへ送る。

2 問題文の重要な条件

- S3 Event Notificationsはat-least-onceで、重複と順序逆転が起こり得る
- 同じオブジェクトキーの新旧だけを判定し、異なるキーは並列処理したい
- 一時障害時の耐久バッファと再試行に加え、更新側を冪等にする必要がある

3 正解へ至る判断過程

1. 同時実行数やキューの配送順だけでは、発生源からの重複配送と既に逆転した到着順を正しい版順へ戻せない。
2. S3イベントのsequencerは同じオブジェクトキーに対するPUT/DELETEイベントの相対順序を比較するために使えるが、異なるキー間で比較しない。
3. SQSで障害時の再試行を保持し、DynamoDBの条件式でキーごとの最大sequencerとイベントIDを原子的に更新して、処理そのものを冪等化する。

4 正解が要件を満たす理由

DはSQS Standardを耐久バッファとしてLambdaの一時障害やスロットリングを吸収し、DLQで繰返し失敗も隔離します。S3イベントに含まれるsequencerを同じbucket/keyについて比較し、DynamoDBの条件付き書込みで保存済み値より新しいイベントだけを受理すれば、遅れて届いた旧版が最新版を上書きしません。同一イベントの処理IDも条件に含めて重複を無害化し、パーティションキーをbucket/keyにすることで異なるオブジェクトは並列処理できます。

5・6 各選択肢の検討

不適切な理由と、正解になり得る条件を対にして確認します。

A 不正解

同時実行数1は処理を遅くするだけで、S3のat-least-once配送による重複を消さず、通知の到着順も保証しません。全キーの並列性も失われます。

この選択肢が正解になる条件

イベント源が順序と一意配送を別途保証し、処理量が非常に小さく、共有資源保護のため全処理を直列化すること自体が目的なら同時実行数制限が使えます。

B 不正解

単一MessageGroupIdは全キーを直列化します。またFIFOの送信重複排除には時間窓があり、S3由来イベントの永続的な業務冪等性や、既に逆転した版の意味的判定を代替しません。

この選択肢が正解になる条件

イベント発行側が安定した重複排除IDとオブジェクト別MessageGroupIdを設定でき、5分窓内の送信重複だけが問題で、消費処理も冪等ならFIFOが順序制御に適します。

C 不正解

アーカイブは監査や再処理に有効ですが、再生順をオブジェクトの版順として扱えず、無条件更新は重複・旧版による上書きを再発させます。秒単位処理要件にも合いません。

この選択肢が正解になる条件

目的が過去イベントの監査、障害修復、または新しい処理ロジックへの一括再投入で、コンシューマーが独自に冪等性と順序を処理するならEventBridgeアーカイブが有効です。

D 正解

配送の耐久性と再試行をSQSで担い、キーごとのsequencer比較と条件付き書込みで重複・旧版を無害化しながら並列性を保ちます。

7 今後使える判断基準

at-least-onceのイベントでは、配送サービスをexactly-onceだと仮定せず、耐久キューと消費側冪等性を組み合わせる。S3のsequencerは同じオブジェクトキー内だけで比較し、条件付き書込みで状態遷移を原子的に守る。

8 関連サービスとの比較

SQS Standardは高スループットの耐久バッファですが重複・順序逆転を前提とします。SQS FIFOはグループ内順序と送信重複排除を支援しますが、業務上の永続的冪等性は別途必要です。DynamoDBの条件式は処理済みIDや最大sequencerを原子的に保護します。

QUESTION 086

問題86

PERFORMANCE

本番より難しい

単一選択

RDSの高可用性と読取り拡張の分担

旅行予約会社は、単一AZのRDS for PostgreSQL DBインスタンスを使用している。負荷の80%は検索・レポートの読取りで、繁忙期にはCPUが90%を超える一方、書込みはCPUの15%程度である。検索結果は数秒の遅延を許容できるが、予約確定直後の確認は最新データを読む必要がある。またAZ障害や計画メンテナンス時には手動復旧をせず自動フェイルオーバーし、データ損失を最小化したい。採用中のエンジン版では、従来型のRDS Multi-AZ DBインスタンス配置を使用する。アプリケーションは読取り接続先と書込み接続先を分離でき、レプリカ遅延を監視して検索をプライマリへ戻す実装も可能である。コストは増えても単一の最大インスタンスだけに依存したくない。可用性と読取り性能を両方改善する最適な構成はどれですか。

- A** DBをMulti-AZ DBインスタンス配置へ変更し、別途リードレプリカを複数作成する。書込みと整合性必須読取りはプライマリ、遅延許容読取りは各レプリカへ送る。
- B** DBをMulti-AZ DBインスタンス配置へ変更し、同期スタンバイへ検索クエリを直接送る。プライマリ障害時は同じスタンバイを読取りと書込みに同時利用できると見なす。
- C** Single-AZのままクロスAZリードレプリカだけを作成し、すべての読取りを移す。障害時はRDSがそのレプリカを同期スタンバイとして自動的に昇格すると見なす。
- D** Single-AZのDBを最大クラスへスケールアップし、読取りと書込みを同じエンドポイントへ送り続ける。毎日の自動スナップショットからの復元を可用性対策とする。

ここで答えを決める

正解・解説は次のページから

ANSWER 086

問題86の正解・解説

1 正解 **A**

- A** DBをMulti-AZ DBインスタンス配置へ変更し、別途リードレプリカを複数作成する。書込みと整合性必須読取りはプライマリ、遅延許容読取りは各レプリカへ送る。

2 問題文の重要な条件

- 従来型Multi-AZ DBインスタンスの同期スタンバイは高可用性用で読取り先ではない
- 読取り80%のCPU負荷を横に逃がしつつ、数秒のレプリカ遅延を許容できる検索がある
- 書込みと最新読取りにはプライマリを使い、AZ障害には自動フェイルオーバーが必要

3 正解へ至る判断過程

1. 高可用性要件には同期スタンバイへ自動フェイルオーバーするMulti-AZ DBインスタンス配置を選ぶ。
2. この配置のスタンバイは通常の読取り処理には使えないため、CPU負荷軽減には非同期リードレプリカを別途追加する。
3. レプリケーション遅延を考慮し、古さを許容する検索だけをレプリカへ送り、read-after-writeが必要な確認はプライマリへ残す。

4 正解が要件を満たす理由

Aは同期レプリケーションを使うMulti-AZ DBインスタンス配置でプライマリ障害時の自動フェイルオーバーを提供し、可用性を担当させます。さらに非同期リードレプリカへ検索・レポートを分散して、読取りが占めるCPU負荷をスケールアウトします。アプリケーションが接続先を分けられ、検索は数秒の遅延を許容するため、レプリカラグも条件内です。予約確定直後の確認をプライマリへ向ければ、最新性の必要な処理も守れます。

5・6 各選択肢の検討

不適切な理由と、正解になり得る条件を対にして確認します。

A 正解

Multi-AZを可用性、リードレプリカを読取り性能に使い分け、整合性要件に応じて接続先も分離しています。

B 不正解

従来型Multi-AZ DBインスタンスの同期スタンバイは通常アクセスできず、検索クエリのオフロード先として使えません。可用性を読取り容量と二重計上しています。

この選択肢が正解になる条件

RDS Multi-AZ DBクラスターなど、明示的に読み取り可能なスタンバイを持つ別の配置方式を採用でき、そのエンジン・バージョン・費用が要件に合うなら読取りへ利用できます。

C 不正解

リードレプリカは読取り負荷を逃がせますが非同期であり、従来型Multi-AZの同期スタンバイや自動フェイルオーバーの代替ではありません。昇格には別設計が必要です。

この選択肢が正解になる条件

高可用性や自動切替が不要で、目的が遅延許容読取りのスケールだけ、またはDR手順として手動昇格を許容するならリードレプリカ単独が適します。

D 不正解

スケールアップは一時的にCPU余力を増やしますが単一AZ障害点を残し、スナップショット復元は自動フェイルオーバーではありません。読取り負荷の水平分散もできません。

この選択肢が正解になる条件

負荷が短期間だけ増え、利用可能な最大クラス内で収まり、Single-AZ停止とバックアップからの長い復元を許容する非重要環境なら単純なスケールアップも選べます。

7 今後使える判断基準

RDSではMulti-AZとリードレプリカを同じ『複製』として扱わない。従来型Multi-AZは同期スタンバイと自動フェイルオーバーによる可用性、リードレプリカは非同期で遅延を許容する読取りスケールが主目的である。

8 関連サービスとの比較

Multi-AZ DBインスタンスのスタンバイは同期複製され、自動フェイルオーバーに使われますが通常は読取り不可です。リードレプリカは非同期で読取り可能ですがラグがあり、HAの直接代替ではありません。なおMulti-AZ DBクラスターは読取り可能なインスタンスを持つため、配置タイプを問題文で見分ける必要があります。

QUESTION 087

問題87

COST

複合難問

単一選択

小容量オブジェクトの集約と分析費削減

IoT企業は3億台分のセンサーレコードを、1件約2KBのJSONオブジェクトとしてLambdaからS3へ1件ずつPUTしている。1日あたり約4億オブジェクトが作られ、保存容量よりPUT、LIST、GETリクエスト費用が大きく、Athenaの日次集計も大量の小ファイルとJSON全列スキャンで遅い。データは追記専用で、個々のレコードを到着直後にオブジェクトとして取得する要件はなく、15分以内にS3へ永続化されればよい。分析は主に日付、時間、デバイス種別で絞り込み、90日間は随時参照する。送信側プロトコルの変更は可能だが、サーバークラスターの運用は避けたい。障害時の未配送データも確実に再試行できる必要がある。リクエスト総数、総保存量、Athenaスキャン量をまとめて削減する最適な設計はどれですか。

- A** 2KBオブジェクトの書込み方法は維持し、S3 Intelligent-Tieringへ直接保存する。アクセス階層の自動移動だけでPUT数とAthenaのファイルオープン数も削減する。
- B** レコードをAmazon Data Firehoseへ送り、時間とサイズでバッファして大きなS3オブジェクトに集約する。日付・時間で分割したParquetへ変換し、圧縮して保存する。
- C** 2KBオブジェクトをS3 Standardへ個別PUTし、1日後にGlacier Deep Archiveへ移行する。Athena集計のたびに対象オブジェクトを一括復元する。
- D** 全オブジェクトをS3 Express One Zoneへ個別PUTし、Athenaから直接読み取る。単一AZの低レイテンシによりリクエスト数とスキャンバイトも自動的に減ると見なす。

ここで答えを決める

正解・解説は次のページから

ANSWER 087

問題87の正解・解説

1 正解 **B**

- B** レコードをAmazon Data Firehoseへ送り、時間とサイズでバッファして大きなS3オブジェクトに集約する。日付・時間で分割したParquetへ変換し、圧縮して保存する。

2 問題文の重要な条件

- 2KBオブジェクトを1日4億件作るため、容量よりリクエスト数と小ファイルが支配的
- レコード単位の即時取得は不要で、最大15分のバッファを許容できる
- Athenaは時間・種別で絞るため、集約、列指向形式、圧縮、パーティションを同時に活用できる

3 正解へ至る判断過程

1. ストレージクラス移行は保存単価を変えても、最初の4億PUTとAthenaが開く小ファイル数を減らさない。
2. 許容遅延内でストリームをバッファし、複数レコードを1オブジェクトへまとめれば、S3へのPUTと後続GET/LISTの回数を桁違いに減らせる。
3. JSONを圧縮Parquetへ変換し、よく使う日付・時間で過剰にならない粒度に分割すれば、Athenaが読む列とパーティションを削減できる。

4 正解が要件を満たす理由

BはマネージドなFirehoseのバッファサイズ・間隔を使い、数万件以上の小レコードをより大きなS3オブジェクトへまとめるため、PUT数と小ファイル処理を大幅に減らします。Parquetは必要列だけを読みやすい列指向形式で、圧縮と日付・時間パーティションを組み合わせれば保存バイトとAthenaのスキャンバイトも削減できます。15分以内という永続化要件はFirehoseのバッファ設定内で満たせ、集約用EC2クラスターのパッチやスケール運用も不要です。

5・6 各選択肢の検討

不適切な理由と、正解になり得る条件を対にして確認します。

A 不正解

Intelligent-Tieringはアクセス頻度に応じた保存階層を最適化しますが、オブジェクト数、最初のPUT回数、小ファイルを開く分析処理、JSONの全列スキャンを減らしません。

この選択肢が正解になる条件

オブジェクトが十分大きく、アクセス時期を予測できず、主な費用が保存容量であってリクエスト・分析スキャンが問題でないならIntelligent-Tieringが適します。

B 正解

許容された遅延を使って小レコードを集約し、Parquet・圧縮・適切なパーティションでリクエスト、保存、Athenaの三つの費用要因を同時に下げます。

C 不正解

個別PUT費用は残り、90日間随時分析するデータをDeep Archiveへ早期移行すると復元待ち、復元費、最低保存期間の影響が加わります。小ファイル問題も解消しません。

この選択肢が正解になる条件

年単位でほぼ再読せず、復元に時間がかかってよく、既に大きなオブジェクトへ集約済みの長期アーカイブならDeep Archiveが適します。

D 不正解

Express One Zoneは低レイテンシ・高性能アクセス向けですが、4億回の個別PUTやJSONスキャン量そのものは減らしません。単一AZという耐障害性のトレードオフも持ち込みます。

この選択肢が正解になる条件

単一AZの高性能コンピュータから同じデータへ非常に高頻度・低レイテンシでアクセスすることが主目的で、データを別途保護できるならExpress One Zoneが候補です。

7 今後使える判断基準

S3費用を容量だけで見ず、オブジェクト平均サイズ、PUT/GET/LIST件数、分析時スキャン量を分ける。大量の小レコードで即時個別取得が不要なら、保存前に集約し、圧縮列指向形式とクエリ述語に合う粗すぎないパーティションを選ぶ。

8 関連サービスとの比較

Data Firehoseはストリームのバッファ、S3配送、形式変換、圧縮をマネージドに行います。Intelligent-TieringやGlacierは主に保存階層を最適化し、S3 Express One Zoneはアクセス性能を最適化します。いずれも単独では小オブジェクトのリクエスト数を集約しません。

QUESTION 088

問題88

SECURE

本番相当

複数選択 (2つ)

専有HSMとKMS統合を両立する鍵ストア

決済会社はS3、EBS、RDSの暗号化にKMS APIと各サービスのSSE-KMS統合を使い続けたい。一方、新しい監督規則では、鍵素材を自社AWSアカウント内の専有シングルテナントHSMで生成・保持し、暗号処理中もそのHSM外へ平文で出さないこと、HSMクラスターとバックアップを自社が管理できることが求められる。対象は対称暗号鍵であり、オンプレミスに鍵を残す要件はない。1 AZ障害で暗号化処理が全面停止しない構成とし、一般的なKMS対応AWSサービスのコード変更を最小化したい。標準KMSキーが使うAWS管理のHSMは認証要件を満たすが、専有性の要件を満たさないと監査人が明記している。追加のHSM運用費は既に承認済みで、可用性を顧客側でも設計する方針である。実施すべき対応を2つ選択してください。

- A** 通常のカスタマーマネージドKMSキーを作成し、キーポリシーで会社だけに利用を限定する。AWS KMS内の論理的なアクセス分離を専有シングルテナントHSMと同等として扱う。
- B** AWS CloudHSMクラスターを作成し、少なくとも2台のアクティブHSMを異なるAZへ配置する。HSMユーザー、クラスター、バックアップの管理手順を整備する。
- C** CloudHSMクラスターを基盤とするAWS KMSのCloudHSMカスタムキーストアを作成・接続し、その中に対称暗号KMSキーを作って対象AWSサービスから使用する。
- D** オンプレミスHSMの鍵をAWS KMS External Key Storeへ接続し、すべての暗号処理を社内設備で行う。ネットワーク停止時の可用性はAWS KMSが自動補完すると見なす。

ここで答えを決める

正解・解説は次のページから

ANSWER 088

問題88の正解・解説

1 正解 **B・C**

- B** AWS CloudHSMクラスターを作成し、少なくとも2台のアクティブHSMを異なるAZへ配置する。HSMユーザー、クラスター、バックアップの管理手順を整備する。
- C** CloudHSMクラスターを基盤とするAWS KMSのCloudHSMカスタムキーストアを作成・接続し、その中に対称暗号KMSキーを作って対象AWSサービスから使用する。

2 問題文の重要な条件

- 鍵素材を自社アカウント内の専有シングルテナントHSMに置く必要がある
- 一般的なAWS KMS APIとAWSサービス統合を維持したい
- 1 AZ障害を考慮してCloudHSMとカスタムキーストアを高可用にする

3 正解へ至る判断過程

1. 通常のKMSキーは強い論理分離と認証済みHSMを提供するが、監査が要求する顧客専有CloudHSMクラスターではない。
2. CloudHSMだけを直接使うとPKCS#11等のアプリケーション統合になり、S3・EBS・RDSのKMS統合をそのまま利用する目的に足りない。
3. 複数AZのCloudHSMクラスターにKMS CloudHSMカスタムキーストアを接続することで、専有HSMの統制とKMSインターフェイスを組み合わせる。

4 正解が要件を満たす理由

Bで異なるAZに複数のアクティブHSMを持つ専有CloudHSMクラスターを用意し、単一HSM・単一AZ障害を避けながら、会社がHSMユーザーとバックアップを管理します。CでそのクラスターをKMSのCloudHSMカスタムキーストアへ関連付けると、KMSが非抽出の対称鍵素材を専有HSM内に生成し、通常のKMS APIと対応AWSサービス統合を通じて利用できます。両方を組み合わせて初めて、専有性、高可用性、顧客統制、既存統合の条件を同時に満たします。

5・6 各選択肢の検討

不適切な理由と、正解になり得る条件を対にして確認します。

A 不正解

標準KMSキーはキーポリシーによる論理分離と強固なHSM保護を提供しますが、設問で明示された顧客専有シングルテナントHSMという物理的専有要件を満たしません。

この選択肢が正解になる条件

規制がFIPS認証済みHSM、鍵の非平文保護、顧客管理ポリシーだけを求め、専有HSMやHSM運用統制を要求しないなら標準KMSキーが最も低運用です。

B 正解

顧客が管理するシングルテナントHSMを複数AZへ配置し、カスタムキーストアが必要とする可用な基盤を提供します。

C 正解

専有CloudHSM内の鍵素材を、KMS APIとS3・EBS・RDSなどの統合から利用できるようにする橋渡しです。

D 不正解

External Key Storeは鍵をAWS外部に保持する要件向けで、今回は自社AWSアカウント内CloudHSMが指定されています。また外部鍵管理システムとネットワークの可用性・遅延は顧客責任で、KMSが自動補完しません。

この選択肢が正解になる条件

法規が鍵をAWS外またはオンプレミスでのみ保持・使用することを要求し、外部HSM、XKS Proxy、ネットワークを必要な可用性と低遅延で自社運用できるならExternal Key Storeが適します。

7 今後使える判断基準

暗号サービス選択では『認証レベル』『鍵の所在』『HSMのテナンシー』『誰が可用性を運用するか』『AWSサービス統合』を分ける。標準KMSで専有性が不足し、AWS内専有HSMとKMS統合の両方が必要ならCloudHSMカスタムキーストアを選ぶ。

8 関連サービスとの比較

標準AWS KMSは最も運用が少なく広いAWS統合を持ちます。CloudHSMは顧客管理のシングルテナントHSMと低レベル暗号APIを提供し、KMS CloudHSMカスタムキーストアは両者を接続します。External Key Storeは鍵と暗号処理をAWS外に置く別の信頼モデルです。

QUESTION 089

問題89

RESILIENT

本番より難しい

単一選択

アプリケーション異常を検知した自動置換

動画配信会社は、3つのAZにまたがるApplication Load Balancer（ALB）と、希望容量6台のEC2 Auto Scalingグループで会員APIを実行している。最近、Javaプロセスがデッドロックして依存DBへ接続できなくなる障害が発生した。OSとネットワークは正常なためEC2ステータスチェックには合格するが、アプリケーションの`/ready`は503を返す。ALBはそのインスタンスへの転送を停止するものの、Auto Scalingグループはインスタンスを置換せず、障害が重なり処理可能台数が減る。新規インスタンスは起動後、キャッシュ構築に最大7分かかる。運用チームは独自の監視Lambdaや手動終了を保守せず、ALBが検出したアプリケーション異常に基づいて希望容量を自動回復したい。一時的な起動中エラーによる置換ループも避ける必要がある。最も適切な構成はどれか。

- A** ALBターゲットグループのヘルスチェックパスを`/ready`へ変更し、Unhealthy thresholdを小さくする。Auto Scalingグループのヘルスチェック種別は既定のEC2のまま維持する。
- B** EC2詳細モニタリングを有効にし、StatusCheckFailedメトリクスのアラームでターゲット追跡スケリングを行う。ALBのヘルスチェック結果はルーティング制御だけに使用する。
- C** ALBターゲットグループをAuto Scalingグループへ関連付け、追加のヘルスチェック種別としてELBを有効にする。`/ready`を検査し、起動時間を上回るヘルスチェック猶予期間を設定する。
- D** Route 53でALBのエンドポイントに対するヘルスチェックを作成し、異常時に同じリージョンの別名ALBへフェイルオーバーする。各Auto ScalingグループはEC2ヘルスチェックだけを使う。

ここで答えを決める

正解・解説は次のページから

ANSWER 089

問題89の正解・解説

1 正解 C

- C** ALBターゲットグループをAuto Scalingグループへ関連付け、追加のヘルスチェック種別としてELBを有効にする。`/ready`を検査し、起動時間を上回るヘルスチェック猶予期間を設定する。

2 問題文の重要な条件

- アプリケーションは503を返すがEC2ステータスチェックには合格する
- ALBは異常ターゲットを除外しているが、希望容量を回復する置換が起きていない
- 起動後最大7分の準備時間を許容しつつ、独自監視コードなしで自動復旧したい

3 正解へ至る判断過程

1. ALBのターゲットヘルスチェックは転送可否を判断できるが、その結果をAuto Scalingが利用しなければインスタンス置換には結び付かない。
2. Auto Scalingでは既定でEC2の状態を確認するため、ELBヘルスチェックを追加してALBがUnhealthyとしたインスタンスをAuto Scalingにも異常として扱わせる必要がある。
3. アプリケーション検査を`/ready`に合わせ、7分超のヘルスチェック猶予期間を設定すれば、実障害を置換しつつ初期化中の誤判定を抑えられる。

4 正解が要件を満たす理由

CではALBがアプリケーションレベルで異常を検出すると、その結果を受けたAuto ScalingがインスタンスをUnhealthyと判断して置換し、希望容量6台へ戻す。ターゲットグループの関連付けにより登録も自動化される。さらに起動時間を考慮した猶予期間により、キャッシュ構築中の503で新規インスタンスが連続終了するリスクを抑え、可用性と運用負荷の要件を同時に満たす。

5・6 各選択肢の検討

不適切な理由と、正解になり得る条件を対にして確認します。

A 不正解

ALBは不健全なターゲットへの転送を止められるが、Auto ScalingがELBヘルスチェックを利用しないままでは、EC2自体が正常なインスタンスを置換しない。処理可能台数の減少が残る。

この選択肢が正解になる条件

障害が短時間で自己回復し、異常ターゲットを一時的にルーティングから外すだけでよく、希望容量に対する自動置換が不要なら適切になり得る。

B 不正解

詳細モニタリングとEC2ステータス系メトリクスはOS・基盤障害の観測には有効だが、今回のデッドロックでは値が正常なためスケーリングを起動できない。台数追加と異常個体の置換も同義ではない。

この選択肢が正解になる条件

障害原因がCPU負荷やEC2ステータスチェック失敗として確実に観測され、負荷に応じた容量追加が主目的ならCloudWatchメトリクスによるスケーリングが正解になり得る。

C 正解

ALBのアプリケーションヘルスをAuto Scalingの置換判定へ連携し、猶予期間で起動中の誤置換も防ぐため、全条件を満たす。

D 不正解

Route 53はALB全体やリージョン単位の切替には使えるが、個々の異常EC2を希望容量内で置換しない。別ALBと別グループの常設は費用と運用も増やす。

この選択肢が正解になる条件

独立した2つのスタックまたはリージョンを持ち、インスタンス単位ではなくエンドポイント全体の障害に対してDNSフェイルオーバーする要件なら正解になり得る。

7 今後使える判断基準

ロードバランサーが異常ターゲットを外すことと、Auto Scalingが異常インスタンスを置換することは別である。アプリケーション障害まで自動回復させるなら、ELBヘルスチェックをAuto Scalingへ明示的に連携し、起動時間に合う猶予期間を設定する。

8 関連サービスとの比較

EC2ヘルスチェックはインスタンスや基盤の健全性、ALBターゲットヘルスチェックは指定パスでのアプリケーション応答を確認する。ALB単体はルーティングから除外するだけだが、Auto ScalingでELBヘルスチェックを有効にすると置換まで進む。Route 53ヘルスチェックはスタックやリージョン単位のDNS切替に向く。

QUESTION 090

問題90

PERFORMANCE

複合難問

単一選択

DynamoDB読取りのマイクロ秒化

ECサイトは商品カタログをDynamoDBに保存し、GetItemとQueryを毎秒30万回実行している。アクセスの85%は上位1%の同じ商品に集中し、現在の読取りレイテンシは一桁ミリ秒だが、セール中だけp99を1ミリ秒未満にする必要がある。カタログ表示では数秒の古さを許容でき、既定の結果整合性のある読取りを利用している。商品更新は全体の1%未満で、アプリケーションはDynamoDB SDKのデータプレーンAPIを広く使用している。チェックアウト時の少数の強い整合性を持つ読取りは従来どおりDynamoDBへ到達してよい。開発チームはキャッシュキー設計、キャッシュミス時の読込み、更新時の無効化処理を独自実装したくない。VPC内でMulti-AZの可用性を確保し、キャッシュヒットによってテーブルの読取り消費も減らしたい。最も適切な構成はどれか。

- A** ElastiCache for RedisをMulti-AZで配置し、商品IDをキーとするキャッシュアサイドを実装する。更新処理はDynamoDB書込み後にRedisキーを削除し、ミス時はアプリケーションが再投入する。
- B** ElastiCache for Memcachedのノードを複数AZへ配置し、短いTTLでQuery結果を保存する。各アプリケーションにDynamoDBとMemcachedの二重書込みおよびノード障害時の再構築処理を追加する。
- C** DynamoDBをオンデマンドモードへ変更して最大スループットを引き上げ、すべての読取りを強い整合性へ変更する。キャッシュは使わず、Adaptive Capacityだけで1ミリ秒未満を実現する。
- D** 複数AZにノードを持つDAXクラスターを作成し、アプリケーションのDynamoDBクライアントをDAXクライアントへ置き換える。結果整合性のある読取りと書込みをDAX経由にする。

ここで答えを決める

正解・解説は次のページから

ANSWER 090

問題90の正解・解説

1 正解 **D**

- D** 複数AZにノードを持つDAXクラスターを作成し、アプリケーションのDynamoDBクライアントをDAXクライアントへ置き換える。結果整合性のある読取りと書き込みをDAX経由にする。

2 問題文の重要な条件

- 同じDynamoDB項目への結果整合性のある読取りが大半で、p99を1ミリ秒未満にしたい
- DynamoDB APIを広く利用しており、独自のキャッシュアサイドや無効化を避けたい
- VPC内でMulti-AZの可用性を確保し、DynamoDBへ届く読取りと消費量を減らしたい

3 正解へ至る判断過程

1. 単に容量を増やすのではなく、反復するホットデータの結果整合性読取りをインメモリで処理してマイクロ秒級へ短縮する必要がある。
2. DAXはDynamoDB互換APIのリードスルー／ライトスルーキャッシュであり、キャッシュミス処理や書き込み後の項目キャッシュ更新をサービス側で扱う。
3. 強い整合性の読取りはDAXでキャッシュされずDynamoDBへパススルーされるため、少数のチェックアウト処理を維持しながらカタログ読取りだけを高速化できる。

4 正解が要件を満たす理由

DはDynamoDB向けに設計されたDAXを使い、頻繁に再読込みされる結果整合性データをメモリから返すことで、ミリ秒からマイクロ秒級への短縮とテーブル読取り負荷の削減を狙える。DAXクライアントへの変更で既存APIモデルを保ち、ライトスルー動作により独自の二重書き込みや無効化を減らせる。複数AZにノードを置けばキャッシュ層の可用性も満たす。

5・6 各選択肢の検討

不適切な理由と、正解になり得る条件を対にして確認します。

A 不正解

Redisは汎用キャッシュとして有力で可用性も構成できるが、キャッシュキー、ミス処理、TTL、DynamoDB更新との無効化整合性をアプリケーションが実装する必要があり、変更最小化の条件に劣る。

この選択肢が正解になる条件

複雑なデータ構造、ランキング、分散ロックなどDynamoDB読取り以外にも同じキャッシュを使い、開発チームがキャッシュアサイドと無効化を管理できるなら正解になり得る。

B 不正解

Memcachedは単純な分散キャッシュには適するが、二重書込みや再構築を含むアプリケーション変更が大きく、レプリケーションによる高可用性やDynamoDB専用のライトスルー統合も提供しない。

この選択肢が正解になる条件

失ってもすぐ再生成できる単純な一時キャッシュで、永続性やレプリケーションを求めず、最小機能と水平分散を優先するなら正解になり得る。

C 不正解

オンデマンドモードやAdaptive Capacityはスループット不足を軽減するが、DynamoDBの通常読取りをキャッシュのマイクロ秒レイテンシへ変えるものではない。強い整合性化は消費量も増やす。

この選択肢が正解になる条件

問題がレイテンシではなく予測不能な容量超過によるスロットリングで、単一項目の通常ミリ秒レイテンシを許容できるならオンデマンドモードが正解になり得る。

D 正解

DynamoDB互換の管理されたキャッシュで、結果整合性読取りのマイクロ秒化、読取り削減、実装負荷の低減を同時に満たす。

7 今後使える判断基準

DynamoDBの同じキーやQuery結果を反復して読み、結果整合性を許容し、コード変更を小さくしたいならDAXを優先する。強い整合性読取りはDAXキャッシュを利用しないため、整合性要件と高速化対象を分けて判断する。

8 関連サービスとの比較

DAXはDynamoDB専用でAPI互換、リードスルー／ライトスルーによりキャッシュ管理を減らす。ElastiCache for Redisは多様なデータ構造や用途に強いが通常はキャッシュアサイドを設計する。Memcachedは単純な揮発キャッシュとして軽量である。DynamoDB On-Demandは容量管理を減らす但しキャッシュによるレイテンシ短縮とは目的が異なる。

QUESTION 091

問題91

COST

本番相当

単一選択

定常大容量ハイブリッド回線の費用

映像制作会社は、東京リージョンのS3とEC2からオンプレミス編集拠点へ、毎月約180TBをほぼ一定の速度でダウンロードしている。必要帯域は平常時600Mbps、ピーク時850Mbpsで、現在は冗長なSite-to-Site VPNを通すが、インターネット混雑で転送時間が延びる。AWS Pricing Calculatorによる対象ロケーションの試算では、1Gbps Direct Connectのポート時間、通信事業者のクロスコネクト固定費、Direct Connectのデータ転送OUTを合計しても、VPN接続時間とインターネット経由のデータ転送OUTより月額が低い。会社はDirect Connectロケーションへ接続でき、導入に2か月かけられる。数時間に一度の短い回線障害では転送を再開でき、完全な二重専用線の費用までは許容しない。転送量が少ない災害時の予備経路は必要である。最も費用対効果の高い構成はどれか。

- A** 1GbpsのDirect Connect接続を主経路にし、ポート時間、クロスコネクト、Direct Connectデータ転送OUTを総額で評価する。既存Site-to-Site VPNの一方を待機系として残す。
- B** 既存の2本のSite-to-Site VPNを維持し、両方をAccelerated VPNへ変更して帯域を束ねる。Direct Connectの固定ポート費用を避け、全転送をインターネット料金のまま処理する。
- C** 異なる2か所のDirect Connectロケーションに各1Gbpsの専用接続を常時配置し、Direct Connect GatewayとTransit Gatewayを追加する。VPNは廃止して専用線だけで冗長化する。
- D** Site-to-Site VPNを1本に減らし、S3 Transfer AccelerationとDataSyncを全転送に利用する。エッジ最適化とDataSyncの従量料金でデータ転送OUT料金を置き換える。

ここで答えを決める

正解・解説は次のページから

ANSWER 091

問題91の正解・解説

1 正解 **A**

- A** 1GbpsのDirect Connect接続を主経路にし、ポート時間、クロスコネクト、Direct Connectデータ転送OUTを総額で評価する。既存Site-to-Site VPNの一方を待機系として残す。

2 問題文の重要な条件

- AWSからオンプレミスへ毎月180TBという大きく安定した送信量があり、1Gbps回線の実効帯域内に収まる
- 試算上はDirect Connectの固定費を含めてもデータ転送総額がVPN経由より安い
- 二重専用線は過剰だが、低利用の障害時予備経路は残す必要がある

3 正解へ至る判断過程

1. Direct Connectは利用量ゼロでも発生するポート時間と事業者費用があるため、回線単価だけでなく対象ロケーションの全費用を比較する必要がある。
2. 本問では定常180TBによりDirect Connectのデータ転送OUT単価差が固定費を上回ることが試算済みで、性能の予測可能性も改善する。
3. 可用性要件は完全な専用線二重化ではなく予備経路で足りるため、主回線をDirect Connect、低使用の既存VPNをバックアップにする構成が最も費用対効果が高い。

4 正解が要件を満たす理由

Aは大容量の定常トラフィックを、試算上総額が安いDirect Connectへ移し、インターネット混雑の影響を減らす。ポート時間とクロスコネクトを明示的に含めるため見かけのデータ単価だけの比較にもならない。平常時にほとんど転送しないVPNを待機系として残すことで、二重の専用接続より固定費を抑えながら障害時経路を確保できる。

5・6 各選択肢の検討

不適切な理由と、正解になり得る条件を対にして確認します。

A 正解

高い定常転送量でDirect Connectの固定費を回収し、安価な待機VPNで必要十分な回線冗長性を得る。

B 不正解

Accelerated VPNには通常のVPN接続時間に加えてGlobal Accelerator関連料金が発生し、データ転送OUTも残る。試算でDirect Connectが安い条件ではコストを増やし、専用接続の予測可能性も得られない。

この選択肢が正解になる条件

短期間でDirect Connectを敷設できず、転送量が小さいか一時的で固定ポート費用を回収できず、インターネット経路の変動を許容するならVPNが正解になり得る。

C 不正解

別口ケーションの二重Direct Connectは高い可用性に有効だが、2本分のポート、事業者回線、必要に応じたTransit Gateway費用が加わる。本問の許容停止時間には過剰である。

この選択肢が正解になる条件

専用経路の障害でも転送停止をほぼ許容せず、異なる設備を使う高可用性接続が必須で、その固定費を正当化できるなら正解になり得る。

D 不正解

S3 Transfer AccelerationとDataSyncは転送方式や運用を改善できるが、オンプレミス向け通信のデータ転送OUTを無料化しない。追加の従量料金が生じ、Direct Connectの費用優位を置き換えない。

この選択肢が正解になる条件

世界各地からS3へインターネット経由でアップロードする利用者の速度改善、または差分検出を伴う管理されたファイル移行が主目的なら各サービスが正解になり得る。

7 今後使える判断基準

Direct ConnectとVPNの費用比較では、Direct Connectのポート時間・事業者費用・データ転送OUTと、VPN接続時間・付加機能・通常のデータ転送OUTを月間転送量で総額比較する。定常大容量ほど専用接続、少量・短期ほどVPNが有利になりやすい。

8 関連サービスとの比較

Direct Connectは固定のポート時間料金を持つ一方、大容量の送信では対象ロケーションのデータ転送OUT料金と予測可能な経路が利点になる。Site-to-Site VPNは迅速で固定費が小さいがインターネット品質と通常の転送料に依存する。二重Direct Connectは可用性を上げるが固定費も増え、Accelerated VPNは追加料金で経路を最適化する。

QUESTION 092

問題92

SECURE

本番より難しい

複数選択 (2つ)

ランディングゾーンの予防と自動是正

金融グループはAWS Control Towerで150アカウントを管理し、開発OUには毎月新しいアカウントが追加される。監査部門は、メンバーアカウントがOrganizationsを離脱する操作と、基盤の監査ログ設定を無効化する操作を、管理者権限を持つ利用者にも実行させたくない。一方、既存および新規リソースについて、暗号化されていないEBSボリュームや全IPv4からSSHを許可するセキュリティグループを継続評価し、違反を中央で確認したい。影響が明確なセキュリティグループ規則は15分以内に自動削除するが、データ破壊につながり得るEBS是正は通知だけにする。アカウントごとの手動IAM設定や定期Lambda走査は避け、OU単位で一貫して展開する必要がある。違反履歴は監査時に確認できなければならない。採用すべき対策を2つ選択してください。

- A** Account Factoryが作る各アカウントの管理者ロールへ同じpermissions boundaryを付け、離脱、ログ無効化、未暗号化EBS作成、0.0.0.0/0のSSH許可をすべてIAMだけで拒否する。
- B** 対象OUでControl Towerの予防的コントロールを有効にし、SCPでOrganizations離脱や監査ログ設定の無効化に関するAPI操作を拒否する。必要な例外はOU設計で分離する。
- C** Organizationsのタグポリシーで`Encrypted=true`と`NoPublicSSH=true`を必須値にする。タグ不一致リソースをControl Towerダッシュボードで非準拠とし、タグ変更でネットワーク規則も自動修復する。
- D** Control Towerの検出コントロールまたは組織向けAWS Configルールで構成を継続評価し、中央集約する。公開SSH規則にはSystems Manager Automationの自動修復を関連付け、EBS違反は通知に留める。

ここで答えを決める

正解・解説は次のページから

ANSWER 092

問題92の正解・解説

1 正解 **B・D**

- B** 対象OUでControl Towerの予防的コントロールを有効にし、SCPでOrganizations離脱や監査ログ設定の無効化に関するAPI操作を拒否する。必要な例外はOU設計で分離する。
- D** Control Towerの検出コントロールまたは組織向けAWS Configルールで構成を継続評価し、中央集約する。公開SSH規則にはSystems Manager Automationの自動修復を関連付け、EBS違反は通知に留める。

2 問題文の重要な条件

- 強い権限を持つメンバーアカウント利用者にも特定API操作を実行させない
- 既存リソースを含む構成違反を継続評価し、中央で可視化したい
- 規則ごとに自動修復と通知を使い分け、OUへ一括展開したい

3 正解へ至る判断過程

1. 操作を事前に阻止する予防統制と、作成後の構成を評価する検出統制は役割が異なるため、両方が必要である。
2. Control Towerの予防的コントロールはSCPでOU配下に上限を設定し、メンバーアカウントのIAM許可があっても対象操作を拒否できる。
3. AWS Configベースの検出と修復アクションを組み合わせれば、既存資源を評価し、公開SSHだけを自動是正するというリスク別の運用を実現できる。

4 正解が要件を満たす理由

BはSCPを基盤とする予防的コントロールによって、メンバーアカウント側の管理者権限を超える権限境界をOUへ適用し、禁止操作を実行前に止める。DはAWS Configによる状態評価を既存・新規リソースへ継続適用し、集約された違反情報を提供する。さらにConfigの修復アクションを公開SSHだけに関連付ければ、15分以内の是正とEBSの通知のみという差別化された対応を満たす。

5・6 各選択肢の検討

不適切な理由と、正解になり得る条件を対にして確認します。

A 不正解

permissions boundaryは付与されたIAMプリンシパルの最大権限を制限するが、アカウント内のすべてのロールやルート利用者を自動的に拘束せず、解除権限の管理も必要になる。既存構成の継続評価にもならない。

この選択肢が正解になる条件

単一アカウントで特定の委任ロールだけが対象となり、そのロール作成・境界変更を別の統制で厳格に保護し、組織全体の検出要件がなければ正解になり得る。

B 正解

SCPによるOU単位の予防統制で、メンバーアカウントのIAM管理者にも禁止APIを許可しない。

C 不正解

タグポリシーはタグの標準化には使えるが、EBS暗号化状態やセキュリティグループのCIDRを実体として検査・修復する仕組みではない。タグ値が正しくても設定が安全とは限らない。

この選択肢が正解になる条件

目的がコスト配賦や所有者タグのキー・値・大文字小文字を組織全体で標準化し、技術的なリソース構成の検査が不要なら正解になり得る。

D 正解

AWS Configベースの継続評価、中央可視化、規則ごとに選べる自動修復または通知により、検出・是正要件を満たす。

7 今後使える判断基準

API操作を実行前に止める要件はSCPベースの予防的コントロール、リソースの現在状態を評価する要件はAWS Configベースの検出コントロールで考える。修復は検出とは別に、影響を評価してSSM Automationなどを選択的に関連付ける。

8 関連サービスとの比較

Control Towerの予防的コントロールは主にSCPで権限の上限を定め、検出コントロールはAWS Configルールなどで違反を報告する。プロアクティブコントロールはCloudFormationでのプロビジョニング前評価に向く。permissions boundaryは個別プリンシパルの委任、タグポリシーはタグ標準化であり、組織の予防・構成検出をそのまま代替しない。

QUESTION 093

問題93

RESILIENT

複合難問

単一選択

継続変更を追従する最小停止DB移行

物流会社はオンプレミスの12TB MySQLデータベースをAurora MySQLへ移行する。データベースは24時間稼働し、毎時約80GBの更新がある。1Gbps Direct Connectは利用可能だが、全量コピーには数日かかる見込みで、業務停止は最終切替時の15分以内に制限される。エンジン互換性の確認と対象テーブルの主キー整備は完了しており、移行期間中もソースのバイナリログを必要な期間保持できる。スキーマ変換は不要である。運用チームは移行前に行数・値の検証を実施し、切替時には未適用変更がほぼゼロであることをCloudWatchで確認したい。独自のレプリケーションコードや長期間の二重書込みは避ける。切替後は短期間ソースを読み取り専用で残し、問題がなければ廃止する。最も適切な移行方法はどれか。

- A** 業務を停止してmysqldumpを取得し、S3へアップロード後にAuroraへ復元する。12TBのエクスポートとインポートが完了するまでソースを読み取り専用で維持する。
- B** AWS DMSでフルロードとCDCを行うタスクを作成し、初期ロード中からソース変更を取得する。検証とCDC latencyを監視し、切替時だけ書込みを止めて残差を適用する。
- C** Aurora MySQLからオンプレミスMySQLへの外部レプリカ接続を設定し、Auroraをソース、オンプレミスをレプリカとして数日同期する。切替時にレプリケーション方向を反転する。
- D** Snowball Edgeへデータベースファイルの初回コピーを取得してAuroraのストレージへ直接インポートする。移送中の変更はS3イベント通知で記録し、切替後に手動適用する。

ここで答えを決める

正解・解説は次のページから

ANSWER 093

問題93の正解・解説

1 正解 **B**

- B** AWS DMSでフルロードとCDCを行うタスクを作成し、初期ロード中からソース変更を取得する。検証とCDCLatencyを監視し、切替時だけ書き込みを止めて残差を適用する。

2 問題文の重要な条件

- 12TBの初回移行中も毎時80GBの更新が続き、停止は15分以内である
- MySQLからAurora MySQLへの互換移行で、ログ保持と主キーなどCDC前提を満たせる
- 管理された検証とレプリケーション遅延監視を使い、独自二重書き込みを避けたい

3 正解へ至る判断過程

1. 初回コピーだけでは数日間に発生する変更が欠落するため、フルロードと並行して継続変更を取得する仕組みが必要である。
2. AWS DMSのフルロード+CDCは、既存データをロードしながらソースログから変更を取得し、初回ロード後にターゲットへ追従適用できる。
3. データ検証とCDC遅延を監視し、最終的にソース書き込みを短時間停止して遅延を解消してから接続先を変えることで、停止時間とデータ欠損を抑えられる。

4 正解が要件を満たす理由

Bは大容量の初期ロード中もMySQLの継続変更をCDCで捕捉するため、数日間ソースを稼働させたままAuroraを最新状態へ近づけられる。DMSの検証とCloudWatchメトリクスで移行品質と遅延を判断し、最終切替時だけ書き込みを止めて残りの変更を適用すれば15分以内の停止を狙える。互換エンジンで前提条件も整っており、独自レプリケーションの保守が不要である。

5・6 各選択肢の検討

不適切な理由と、正解になり得る条件を対にして確認します。

A 不正解

論理ダンプは単純で検証しやすいが、12TBの取得・転送・復元中に書き込みを止めるため、15分以内の停止要件を満たさない。オンライン取得だけでは取得後の変更追従が残る。

この選択肢が正解になる条件

データ量が小さく、数時間以上の停止を許容し、一度限りの単純な同種エンジン移行で継続変更の同期が不要なら正解になり得る。

B 正解

フルロードとCDCを並行し、遅延を収束させて短時間で切り替えられるため、全条件を満たす。

C 不正解

提示された方向では空のAuroraがソース、稼働中オンプレミスがレプリカとなり、データをAWSへ移行できない。方向反転を含む自己管理レプリケーションも運用負荷と整合性リスクが高い。

この選択肢が正解になる条件

Auroraからオンプレミスへ読み取りコピーや退避先を作ることが目的で、MySQLネイティブレプリケーションを管理できるなら構成要素になり得る。

D 不正解

Snowball Edgeは回線制約下のファイル転送には有効だが、データベース物理ファイルをAuroraストレージへ直接取り込む方法ではなく、移送中のトランザクション変更をS3イベントで再現できない。

この選択肢が正解になる条件

対象がS3へ取り込める静的ファイル群で、変更が止まっており、ネットワーク転送に長期間かかる大容量オフライン移行ならSnowball Edgeが正解になり得る。

7 今後使える判断基準

停止時間が全量コピー時間より大幅に短いDB移行では、初回フルロードとCDCを組み合わせ、切替直前に書き込み停止、遅延解消、検証、接続先変更を行う。CDCにはソースログ保持、主キー、LOB、DDLなどの前提確認も必要である。

8 関連サービスとの比較

AWS DMSはフルロード、CDC、または両方を管理し、最小停止移行に向く。論理ダンプやスナップショット復元は単純だが停止またはRPOが長くなりやすい。ネイティブレプリケーションは対応エンジン間で有効だが構築・監視を自ら担う。Snowball Edgeは静的データの物理輸送であり、DBのトランザクションCDCとは異なる。

QUESTION 094

問題94

PERFORMANCE

本番相当

単一選択

大量ログの近リアルタイム検索基盤

オンラインゲーム会社は、ECS上の200サービスから1日2TBのJSONログを収集する。運用担当者は発生後60秒以内のログを、自由なメッセージ文字列、ユーザーID、時間範囲の組合せで検索し、ステータスコードやリージョン別に集計してダッシュボード表示したい。典型的な検索は直近3日でp95 5秒以内、4〜30日前のログは多少遅くてもよい。検索利用者は数百人に増える見込みである。現在のRDS for PostgreSQLではJSONBと全文索引が120TBまで増え、ログ取込が業務DBのI/Oを圧迫し、索引保守も負担になっている。原本は監査用S3へ1年間保持できる。業務トランザクションDBをログ検索から分離し、サーバーの手動管理を減らしつつ、検索性能と古いデータのコストを両立したい。最も適切な構成はどれか。

- A** RDS for PostgreSQLを最大クラスへ変更し、ログ専用のリードレプリカを追加する。JSONBと全文索引を全30日分維持し、レプリカで検索とダッシュボード集計を実行する。
- B** 全ログを圧縮ParquetとしてS3へ時間パーティションで保存し、Athenaだけで60秒ごとにダッシュボードを再計算する。索引は作らず、任意の文字列検索も毎回全対象ファイルを走査する。
- C** Data FirehoseなどでAmazon OpenSearch Serviceへログを取り込み、時間単位または日単位のインデックスを作る。Multi-AZ構成とDashboardsを使い、古い読取り専用インデックスはUltraWarmへ移す。
- D** ログをAmazon Redshift Serverlessへ逐次COPYし、SUPER型に全JSONを格納する。全文文字列検索用の索引は使わず、30日分を同じワークグループでスキャンして可視化する。

ここで答えを決める

正解・解説は次のページから

ANSWER 094

問題94の正解・解説

1 正解 C

- C** Data FirehoseなどでAmazon OpenSearch Serviceへログを取り込み、時間単位または日単位のインデックスを作る。Multi-AZ構成とDashboardsを使い、古い読取り専用インデックスはUltraWarmへ移す。

2 問題文の重要な条件

- 1日2TBの半構造化ログを60秒以内に検索可能にする
- 自由な全文検索、フィールド絞込み、集計、ダッシュボードをp95 5秒以内で実行したい
- 業務RDSから分離し、直近3日と4～30日で性能と保存コストを分けたい

3 正解へ至る判断過程

1. ワークロードはトランザクション処理ではなく、大量の時系列ログに対する全文検索とファセット集計であるため、RDBの垂直拡張より検索エンジンへ分離する。
2. OpenSearchは転置インデックスによる全文検索、構造化フィールドのフィルターと集計、Dashboardsによる可視化を同じ基盤で提供する。
3. 時系列インデックスとホット／UltraWarm階層を使えば、直近データの応答性能を保ちつつ古い読取り中心データのノードコストを下げられる。

4 正解が要件を満たす理由

Cはログ分析と全文検索に適したOpenSearchへ業務DBから負荷を分離し、管理された取込とMulti-AZドメインで運用負荷と可用性を改善する。検索対象フィールドをインデックス化することで、メッセージ検索、フィルター、集計を低遅延で実行できる。直近3日はホットデータノード、古い読取り専用インデックスはUltraWarmとするため、30日すべてを高価なホットストレージへ置く必要もない。

5・6 各選択肢の検討

不適切な理由と、正解になり得る条件を対にして確認します。

A 不正解

リードレプリカは検索負荷の一部を分離できるが、主DBからのログ取込、120TBのJSONB索引、レプリケーションI/Oは残る。全文ログ検索をRDBで拡大し続けると費用と索引運用が大きい。

この選択肢が正解になる条件

データ量が小さく、厳密なACID更新や複雑な関係結合が中心で、限定された全文検索だけを既存SQLと一体で行うならRDSが正解になり得る。

B 不正解

S3とAthenaは長期保管やアドホック分析の費用効率に優れるが、索引なしの任意文字列走査を60秒間隔で繰り返してp95 5秒を安定して満たす用途には向かない。

この選択肢が正解になる条件

検索が数分待てる低頻度の監査調査で、日付やリージョンのパーティションを十分に絞れ、走査料金の低さを最優先するなら正解になり得る。

C 正解

近リアルタイム取込、全文検索、集計、可視化、時系列の階層化を一つの管理サービスで実現する。

D 不正解

Redshiftは構造化データの大規模SQL集計には強いが、自由なログメッセージ全文検索を転置索引で処理する選択ではない。30日分のSUPER走査は低遅延検索と費用の両面で不利になる。

この選択肢が正解になる条件

主な処理が定型的な列指向集計、多表結合、BIレポートで、全文検索が不要か前処理で構造化できるならRedshiftが正解になり得る。

7 今後使える判断基準

半構造化ログを近リアルタイムに全文検索・フィルター・集計するならOpenSearchを検討する。ACID業務処理はRDS、低頻度のS3アドホック走査はAthena、構造化された大規模BI集計はRedshiftと、検索形態で分ける。

8 関連サービスとの比較

OpenSearchは転置索引と集計を備え、ログ検索や運用ダッシュボードに適する。RDSはトランザクションと関係整合性、AthenaはS3上のサーバーレスな低頻度SQL、Redshiftは列指向データウェアハウスに強い。OpenSearchのUltraWarmは古い読取り専用インデックスのコストを下げるが、ホット層より検索レイテンシは高くなる。

QUESTION 095

問題95

SECURE

本番より難しい

単一選択

ネットワーク脅威の検出と証跡集約

企業はOrganizations配下の80アカウントを4リージョンで利用している。あるEC2インスタンスから既知の暗号資産マイニング用IPへ大量通信が発生したが、現在は各チームがセキュリティグループとCloudWatchメトリクスを個別確認しており、発見まで2日かかった。セキュリティ部門は、脅威インテリジェンスや異常行動分析を用いて同様の通信と不審なAPI・DNS活動を継続検出したい。また、受信・拒否を含むネットワークフローの原記録を監査アカウントへ1年間保存し、インシデント時に送信元・宛先・ポート・バイト数をAthenaで調査する必要がある。全アカウント・全リージョンの検出結果は単一の管理画面へ集約し、新規アカウントにも自動適用したい。独自SIEM関連ルールの開発は最小化する。最も適切な構成はどれか。

- A** 全VPCでFlow Logsを有効にして監査アカウントのS3へ集約し、Athenaの保存済みクエリを毎日手動実行する。GuardDutyとSecurity Hubは使わず、検出ルールをSQLで管理する。
- B** Organizations統合のGuardDutyを委任管理者から全アカウントへ自動有効化する。GuardDutyの検出結果だけを1年間保持し、VPC Flow Logsの明示的な保存とSecurity Hub集約は行わない。
- C** Security Hubを全アカウントで有効にし、AWS Foundational Security Best Practices標準だけを適用する。GuardDutyは無効のまま、Security HubがVPCトラフィックとDNSを直接解析する。
- D** GuardDutyを委任管理者から組織全体へ自動有効化し、全VPCのFlow Logsを中央S3へ配信する。GuardDuty結果をSecurity Hubへ統合し、委任管理者とクロスリージョン集約を設定する。

ここで答えを決める

正解・解説は次のページから

ANSWER 095

問題95の正解・解説

1 正解 D

- D** GuardDutyを委任管理者から組織全体へ自動有効化し、全VPCのFlow Logsを中央S3へ配信する。GuardDuty結果をSecurity Hubへ統合し、委任管理者とクロスリージョン集約を設定する。

2 問題文の重要な条件

- 既知脅威と異常行動を使い、ネットワーク・API・DNSの不審活動を管理サービスで検出した
- 検出結果だけでなく、ネットワークフローの原記録を中央S3へ1年間保存して調査したい
- 80アカウント・4リージョンと新規アカウントの結果をセキュリティ部門へ一元化したい

3 正解へ至る判断過程

1. VPC Flow LogsはIPトラフィックの証跡を保存するが、脅威インテリジェンスや異常分析による検出エンジンではないためGuardDutyと役割を分ける。
2. GuardDutyは独立したフローログ、CloudTrail管理イベント、Route 53 Resolver DNSログなどを分析して脅威を検出するが、調査用の生Flow Logsを利用者のS3へ保存する代替にはならない。
3. Security Hubの組織統合とクロスリージョン集約でGuardDutyなどの検出結果を中央管理し、別途S3のFlow LogsをAthena調査へ残すDが全要件を満たす。

4 正解が要件を満たす理由

DはGuardDutyの管理された脅威検出、VPC Flow Logsの長期生証跡、Security Hubの所見集約をそれぞれの役割に割り当てる。Organizationsの委任管理者と自動有効化で新規アカウントを取り込み、Security Hubのホームリージョンへリンクリージョンの所見を集約できる。中央S3のFlow LogsはGuardDuty所見の前後をAthenaで詳細調査できるため、検出速度とフォレンジック証跡を両立する。

5・6 各選択肢の検討

不適切な理由と、正解になり得る条件を対にして確認します。

A 不正解

Flow LogsとAthenaは事後調査には有効だが、SQLを毎日実行するだけでは既知脅威、DNS、API異常を継続的に相関検出できず、発見時間とルール保守の要件を満たさない。

この選択肢が正解になる条件

監査証跡の保存と低頻度の既知条件による事後調査だけが必要で、リアルタイムに近い脅威検出やAPI・DNS分析が不要なら正解になり得る。

B 不正解

GuardDutyの組織展開は脅威検出を満たすが、所見はVPC Flow Logsの各通信レコードそのものではない。1年間の原記録とAthenaでの詳細調査、複数リージョンの統合ビューが不足する。

この選択肢が正解になる条件

必要なのがGuardDuty所見の組織集約だけで、ネットワーク原記録の長期保管やSecurity Hubによる他製品所見との統合が不要なら正解になり得る。

C 不正解

Security Hubは標準への準拠状態や統合サービスの所見を集約するが、単独でVPCフローやDNSを脅威インテリジェンスと照合する検出サービスではない。原記録も保存しない。

この選択肢が正解になる条件

GuardDutyなど必要な検出サービスが既に有効で、課題が複数製品の所見とセキュリティ標準の中央集約だけならSecurity Hubの展開が正解になり得る。

D 正解

GuardDutyで検出し、Flow Logsで原記録を残し、Security Hubでアカウント・リージョン横断集約するため、全条件を満たす。

7 今後使える判断基準

VPC Flow Logsは通信の観測・証跡、GuardDutyは管理された脅威検出、Security Hubは所見と準拠状態の集約である。検出、詳細調査、組織全体の一元管理を同時に求める問題では、三者を代替関係ではなく補完関係として組み合わせる。

8 関連サービスとの比較

VPC Flow LogsはENIを通過するIPフロー情報をCloudWatch Logs、S3、Data Firehoseへ配信する。GuardDutyは独立したデータストリームを解析して所見を生成し、Flow Logsの保存設定を必要としない。Security HubはGuardDutyなどの所見を標準形式で集約し、クロスリージョンのホームリージョン表示を提供するが、生トラフィックを直接検出するサービスではない。

QUESTION 096

問題96

RESILIENT

複合難問

複数選択 (2つ)

ステートフルEC2の低コスト地域復旧

製造会社は単一リージョンのEC2でライセンス管理アプリケーションを実行している。ルートEBSにはOSと署名済みアプリケーションを置き、2TBの別EBSに頻繁に変わる状態データを保存する。リリースは月1回で、DRリージョンでは承認済みリリースと同一イメージから起動しなければならない。リージョン障害に対するRPOは4時間、RTOは6時間である。DR側の常時稼働EC2や継続レプリケーション用サーバーは持てず、ホストへの追加エージェントも禁止されている。バックアップはDRリージョンの中央バックアップアカウントへコピーして一元監査し、復旧時は選んだリカバリポイントをワークロードアカウントのDR保管庫へcopy-backしてから復元する。実測では毎時バックアップの中央コピーが3時間以内、copy-back、復元、起動が6時間以内に完了する。復旧用VPC、IAM、起動テンプレートはIaCで準備済みである。採用すべき対策を2つ選択してください。

- A** AWS Elastic Disaster RecoveryのエージェントをEC2へ導入し、DRリージョンの低コストステージング領域へブロックを継続複製する。障害時はリカバリインスタンスを起動する。
- B** Data Lifecycle ManagerでデータEBSのスナップショットを1時間ごとに取得し、同じ本番アカウントのDRリージョンへコピーする。中央AWS Backupポリシーと保管先は使用しない。
- C** 各承認リリースでEBS-backed AMIを作成し、DRリージョンの顧客管理KMSキーで暗号化してコピーする。DR側起動テンプレートが承認AMI IDを参照するようIaCで更新する。
- D** AWS Backupの組織的なプランでデータEBSを1時間ごとに保護し、DRリージョンの中央バックアップアカウントへクロスリージョン・クロスアカウントコピーする。復旧時のワークロードアカウントへのcopy-backと復元まで定期訓練する。

ここで答えを決める

正解・解説は次のページから

ANSWER 096

問題96の正解・解説

1 正解 C・D

- C** 各承認リリースでEBS-backed AMIを作成し、DRリージョンの顧客管理KMSキーで暗号化してコピーする。DR側起動テンプレートが承認AMI IDを参照するようIaCで更新する。
- D** AWS Backupの組織的なプランでデータEBSを1時間ごとに保護し、DRリージョンの中央バックアップアカウントへクロスリージョン・クロスアカウントコピーする。復旧時のワークロードアカウントへのcopy-backと復元まで定期訓練する。

2 問題文の重要な条件

- ・ 承認済みOS・アプリケーションと同一の起動イメージをDRリージョンに事前準備したい
- ・ 毎時取得と3時間以内の中央コピーでRPO 4時間を満たし、別アカウントで一元監査したい
- ・ 常時DRコンピュート、レプリケーションサーバー、ホストエージェントを利用できない

3 正解へ至る判断過程

1. 変更頻度の低い実行環境と、頻繁に変わる状態データを同じ方法・頻度で扱う必要はなく、リリースイメージとデータバックアップに分離する。
2. EBS-backed AMIをリリースごとにクロスリージョンコピーすれば、DRリージョンから承認済み構成をすぐ起動でき、常時EC2を置く必要がない。
3. 状態データはAWS Backupの時間単位計画と別リージョン・別アカウントコピーでRPO、分離、中央監査を満たす。復旧時は中央保管庫からワークロード側DR保管庫へcopy-backして復元する。

4 正解が要件を満たす理由

Cは月次リリースを承認AMIとしてDRリージョンへ前置きし、起動テンプレートから再現可能にする。DはデータEBSを毎時保護して別リージョン・中央アカウントへコピーし、リージョン障害と本番アカウント侵害に備える。通常の中央保管庫から別アカウントへ直接復元できるとは見なせず、災害時のcopy-backをrunbookに含める。AWS Backupのコピー時間にSLAはないため、本問では実測した3時間以内のコピーと6時間以内の全復旧でRPO/RTOを裏付けている。

5・6 各選択肢の検討

不適切な理由と、正解になり得る条件を対にして確認します。

A 不正解

Elastic Disaster Recoveryは短いRPO/RTOの継続ブロック複製に適するが、レプリケーションエージェントとDR側ステージング資源を必要とし、本問の明示的な導入・費用制約に反する。

この選択肢が正解になる条件

数秒から数分のRPOと短いRTOが必要で、ソースへのエージェント導入およびDR側ステージング領域の常時費用を許容できるなら正解になり得る。

B 不正解

DLMのスナップショットとクロスリージョンコピーはEBSのDRに有効だが、指定された中央バックアップアカウントの保管庫、組織ポリシー、統合ジョブ監査を満たさず、承認AMIも準備しない。

この選択肢が正解になる条件

単一アカウントでEBSスナップショットのライフサイクルとクロスリージョンコピーだけを簡潔に自動化し、中央AWS Backupガバナンスや別アカウント分離が不要なら正解になり得る。

C 正解

承認済みルート環境をDRリージョンに再現可能なAMIとして前置きし、RTOと監査要件を満たす。

D 正解

変動データを別リージョン・別アカウントへ保護し、中央監査を実現する。copy-backを含む訓練実測でRPO/RTOも確認している。

7 今後使える判断基準

EC2のDRでは、OS・アプリの再構築物と状態データの復旧物を分ける。中央バックアップアカウントを使う場合は、保存先だけでなく復元先アカウントへのcopy-backをrunbookに含める。バックアップ間隔だけでRPOを決めず、コピー遅延、copy-back、復元、起動を訓練してRPO/RTOを実測する。

8 関連サービスとの比較

AMIコピーは同一構成のEC2を別リージョンで起動するリリース成果物に適する。AWS Backupはスケジュール、保持、クロスリージョン・クロスアカウントコピーを一元管理するが、コピー完了時間は固定保証されず、中央保管庫からの復旧にはcopy-backが必要になり得る。DLMは単一アカウントのEBS自動化、Elastic Disaster Recoveryはエージェントによる継続複製で短いRPOを狙う。

QUESTION 097

問題97

PERFORMANCE

本番相当

単一選択

販売開始バーストの同期API

イベント会社は座席予約APIを新規構築する。平常時は每秒100件だが、販売開始後10分間は每秒3万件まで急増する。クライアントは1秒以内に予約確定または売切れを同期的に受け取る必要があり、同じ冪等性キーの再送で座席を二重確保してはならない。アクセスは十分に分散した予約IDを持ち、販売日の負荷試験では想定ピークまで段階的にテーブルを事前ウォームできる。座席状態は単一の原子的更新で確定できる。サーバー常時運用は避け、ピーク後のアイドル費用を抑えたい。API Gateway、Lambda、DynamoDBのサービスクォータは事前に確認・引上げ可能で、販売直前だけコールドスタート対策費用を許容する。突然の単一パーティション集中を避け、バックエンド容量を超えるリクエストには明示的な再試行応答を返す。最も適切な構成はどれか。

- A** API GatewayからLambdaを同期呼出しし、販売時間帯だけProvisioned Concurrencyを用意する。事前ウォームしたDynamoDB On-Demandへ分散キーで条件付き書込みし、APIスロットリングも設定する。
- B** ALB配下にピーク3万件/秒を常時処理できる固定台数のEC2を配置し、RDS for MySQL Multi-AZで座席行をロックする。販売終了後も容量を維持して再開時の遅延を防ぐ。
- C** API GatewayとLambdaを使い、DynamoDB On-Demandのパーティションキーを全リクエストで同じ販売イベントIDにする。Adaptive Capacityが単一キーへの全書込みを自動分散する。
- D** API GatewayからSQS標準キューへ全予約を送って直ちに202を返し、Lambdaが後でDynamoDBへ書き込む。予約結果は数分後にメールで通知し、同期応答は行わない。

ここで答えを決める

正解・解説は次のページから

ANSWER 097

問題97の正解・解説

1 正解 **A**

- A** API GatewayからLambdaを同期呼出しし、販売時間帯だけProvisioned Concurrencyを用意する。事前ウォームしたDynamoDB On-Demandへ分散キーで条件付き書込みし、APIスロットリングも設定する。

2 問題文の重要な条件

- 平常時と10分間のピークに300倍の差があり、アイドル費用とサーバー運用を抑えたい
- クライアントへ1秒以内に確定結果を返し、再送でも二重予約を防ぎたい
- 負荷試験・事前ウォーム・クォータ引上げが可能で、分散キーを利用できる

3 正解へ至る判断過程

1. 短時間の極端なバーストと低い平常負荷には、リクエスト単位で拡張するAPI GatewayとLambda、容量管理を減らすDynamoDB On-Demandが適する。
2. 同期結果が必要なのでキューへ受理しただけでは足りず、Lambdaから条件付き書込みを完了して結果を返す必要がある。冪等性キーを条件式で一度だけ登録する。
3. On-Demandでも過去ピークを大幅に超える瞬間負荷やホットキーは無条件に吸収しないため、事前ウォーム、分散キー、クォータ、Lambda同時実行とAPIスロットリングまで整えるAを選ぶ。

4 正解が要件を満たす理由

Aは常設サーバーを持たずに短時間ピークへ拡張し、ピーク後のコンピュート費用を下げる。販売時間だけProvisioned Concurrencyを使うことでLambdaの初回遅延を抑え、事前ウォームしたDynamoDB On-Demandと高カーディナリティキーで急増を受ける。条件付き書込みが同一冪等性キーの重複確保を拒否し、同期統合が確定結果を1秒以内に返す。スロットリングは下流保護と再試行制御にもなる。

5・6 各選択肢の検討

不適切な理由と、正解になり得る条件を対にして確認します。

A 正解

バースト拡張、同期応答、冪等性、ホットキー回避、ピーク後の低コストを一体で満たす。

B 不正解

EC2とRDSでも厳密な予約は実装できるが、最大ピークの固定容量を常時維持すると平常時の費用が大きい。接続数、ロック競合、スケール時間の設計も増え、本問の運用制約に劣る。

この選択肢が正解になる条件

負荷が一日中安定し、複雑なSQLトランザクションや既存MySQL互換性が必須で、常設容量の利用率が高いなら正解になり得る。

C 不正解

On-DemandとAdaptive Capacityでも、同じパーティションキーに集中する書込みの物理上限を無くせない。単一イベントIDはホットパーティションを作り、毎秒3万件の書込みを阻害する。

この選択肢が正解になる条件

イベントIDに十分なシャード接尾辞を付けるなど書込みを複数パーティションへ分散し、読取り時の集約設計も受け入れるならサーバーレス構成の一部になり得る。

D 不正解

SQSはバースト吸収と下流保護に優れるが、受理時点では座席確保の成功を確定できず、1秒以内の確定または売切れという同期契約を満たさない。標準キューの重複にも対応が必要である。

この選択肢が正解になる条件

要求を非同期ジョブとして受け付け、202応答と後日の結果通知を利用者が許容し、処理遅延を吸収することが最優先なら正解になり得る。

7 今後使える判断基準

急激なバーストにはサービス名だけでなく、API Gatewayクォータ、Lambda同時実行とコールドスタート、DynamoDBの過去ピーク・ウォームスループット、パーティションキーを一緒に確認する。同期結果が必要ならSQSによる非同期化は契約を変える。

8 関連サービスとの比較

API Gateway+Lambda+DynamoDBは変動の大きいステートレスAPIでアイドル費用を抑えやすい。Provisioned Concurrencyは低遅延を買う代わりに時間料金が発生する。DynamoDB On-Demandは予測不能な容量管理を減らすが、急な過去ピーク超過とホットキーへの設計は必要である。SQSは非同期バッファ、RDSは関係トランザクションに強い。

QUESTION 098

問題98

SECURE

本番より難しい

単一選択

クロスアカウントイベント発行の二層認可

小売企業は、注文アカウントAのECSタスクから、統合アカウントBのEventBridgeカスタムイベントバスへ注文イベントを直接PutEventsで送る。Aには複数のタスクロールと管理者がいるが、`OrderPublisherRole` だけに送信を許可したい。Bは`source` が`com.example.orders`、`detail-type` が`OrderCreated` のイベントだけを受け入れ、他のイベント種別やEventBridgeルールを作成権限は与えない。両アカウントは同じOrganizations配下だが、組織内の別アカウントからの発行も禁止する。監査では送信側と受信側のどちらからでも許可を独立して失効でき、対象イベントバスARNへ最小権限を設定することが求められる。中継Lambdaや長期アクセスキーは使用しない。最も適切な認可構成はどれか。

- A** AのOrderPublisherRoleへBのバスARNに対するevents:PutEventsを許可するIAMポリシーだけを付ける。Bのイベントバスは既定のリソースポリシーのままにする。
- B** BのバスポリシーでAのOrderPublisherRoleにevents:PutEventsだけを許可し、sourceとdetail-typeを条件で限定する。Aの同ロールにもBのバスARNへのPutEventsだけを許可する。
- C** BのバスポリシーでAアカウントのrootプリンシパルにevents:*を許可し、AではOrderPublisherRoleへAdministratorAccessを付ける。イベントパターンで不要なイベントを破棄する。
- D** Bのバスポリシーでaws:PrincipalOrgID条件を使い、組織全アカウントにPutEventsを許可する。B側のルールでsourceとdetail-typeを絞れば発行権限も特定ロールに限定される。

ここで答えを決める

正解・解説は次のページから

ANSWER 098

問題98の正解・解説

1 正解 **B**

- B** BのバスポリシーでAのOrderPublisherRoleにevents:PutEventsだけを許可し、sourceとdetail-typeを条件で限定する。Aの同ロールにもBのバスARNへのPutEventsだけを許可する。

2 問題文の重要な条件

- アカウントAの特定タスクロールだけがアカウントBのカスタムバスへ直接発行する
- PutEvents、対象バス、source、detail-typeまで限定し、ルール管理権限は与えない
- 送信側IAMと受信側バスポリシーから独立して権限を失効できるクロスアカウント認可が必要

3 正解へ至る判断過程

1. クロスアカウントの直接PutEventsでは、受信イベントバスが外部プリンシパルを受け入れるリソースベースポリシーを持つ必要がある。
2. 送信側でも実際に呼び出すOrderPublisherRoleへ、対象バスARNに限定したevents:PutEventsを付与し、他のタスクロールから利用できないようにする。
3. イベントバスポリシーの条件でsourceとdetail-typeを限定すれば、ルーティング後に捨てるのではなく不正な発行そのものを認可段階で拒否できる。

4 正解が要件を満たす理由

Bは受信側のEventBridgeイベントバスリソースポリシーと、送信側ロールのアイデンティティポリシーを両方最小化する。PrincipalをOrderPublisherRole、ActionをPutEvents、Resourceを特定バスに限定し、さらにイベント属性条件を設けるため、組織内の他主体や異なるイベントを発行できない。どちらか一方の許可を削除すれば送信を停止でき、中継処理や静的キーも不要である。

5・6 各選択肢の検討

不適切な理由と、正解になり得る条件を対にして確認します。

A 不正解

送信側IAMが許可しても、BのイベントバスがクロスアカウントプリンシパルのPutEventsをリソースポリシーで許可していないため、アクセスは成立しない。

この選択肢が正解になる条件

送信ロールとイベントバスが同一アカウントにあり、バスの既定アクセス範囲内で、IAMポリシーだけを追加すればよい場合は正解になり得る。

B 正解

送信側IAMと受信側バスポリシーの双方で、ロール、アクション、バス、イベント属性を限定する。

C 不正解

アカウント全体へのevents:*とAdministratorAccessはPutRuleやPutTargetsを含む過剰権限で、特定ロール・イベント種別という条件を満たさない。イベントパターンは発行認可の代替ではない。

この選択肢が正解になる条件

Aの複数管理主体へEventBridgeのバス・ルール・ターゲット管理を委任することが要件で、別途 permissions boundaryやSCPで範囲を管理するなら広い委任が必要な場合がある。

D 不正解

PrincipalOrgIDだけでは同じ組織の全アカウント・許可された多数のプリンシパルが発行可能になり、特定ロール限定に反する。受信ルールはイベントを選別するがPutEvents権限を狭めない。

この選択肢が正解になる条件

組織内の多数アカウントを信頼し、各アカウントから共通バスへの発行を許す集中イベント基盤で、送信側SCPやIAMにより各発行者を別途統制するなら正解になり得る。

7 今後使える判断基準

クロスアカウントのEventBridge発行は、送信プリンシパルのIAM許可と受信バスのリソースポリシーを対で確認する。イベントパターンはルーティング、バスポリシー条件は発行認可であり、sourceやdetail-typeを入口で制限したいときは後者を使う。

8 関連サービスとの比較

EventBridgeイベントバスポリシーは他アカウントからのPutEventsなどを受け入れる入口を定義し、送信側IAMポリシーは誰がその入口を呼べるかを定義する。aws:PrincipalOrgIDは多数の組織アカウントをまとめて信頼する場合に便利だが、単一ロールの最小権限より広い。イベントルールは受信済みイベントの照合と配送を担う。

QUESTION 099

問題99

COST

複合難問

単一選択

SQL Server機能とライセンスを分けた配置

企業は2種類のSQL Serverを6か月以内にAWSへ移行する。基幹系XはEnterprise Editionを24時間使用し、FILESTREAMとホストへ導入する第三者製バックアップエージェントが必須である。既存のコアライセンスはあと3年間利用でき、ライセンス担当者はEC2 Dedicated Hostsへの持込み条件を満たすと確認済みだが、RDSへの持込みに必要な契約条件は満たさない。部門系YはStandard Editionの小規模DBが20個あり、RDSでサポートされる機能だけを使う。Yには既存ライセンスがなく、Multi-AZ、自動バックアップ、パッチ適用の運用削減を求める。XもYも本番のためDeveloper Editionは使えず、エンジン変更を伴う改修は期限内にできない。3年間のライセンス、基盤、運用人件費を含めた総額を最小化しつつ要件を満たす配置はどれか。

- A** XとYをすべてRDS for SQL Server License IncludedのMulti-AZへ移行する。XのFILESTREAMとバックアップエージェントはRDSのマスタユーザー権限で導入する。
- B** XとYをすべてSQL Server License IncludedのEC2へ移行し、各DBを個別インスタンスで運用する。既存Xライセンスは使用せず、バックアップとパッチを各チームが管理する。
- C** XはEC2 Dedicated HostsへBYOLで配置し、必要なHAと運用を自己管理する。Yは互換DBを集約したRDS for SQL Server License IncludedのMulti-AZへ配置する。
- D** XとYをAWS SCTでAurora PostgreSQLへ変換し、DMSで一括移行する。FILESTREAMはS3へ置換し、6か月以内に全アプリケーションのSQLとエージェントを改修する。

ここで答えを決める

正解・解説は次のページから

ANSWER 099

問題99の正解・解説

1 正解 C

- C** XはEC2 Dedicated HostsへBYOLで配置し、必要なHAと運用を自己管理する。Yは互換DBを集約したRDS for SQL Server License IncludedのMulti-AZへ配置する。

2 問題文の重要な条件

- XはRDSで利用できないFILESTREAMとホストエージェントを必要とし、EC2 Dedicated Hostsで使える既存ライセンスがある
- YはRDS対応機能だけで既存ライセンスがなく、Multi-AZ・バックアップ・パッチの管理を減らしたい
- エンジン改修はできず、3年間のライセンス費と運用費を含む総額を比較する

3 正解へ至る判断過程

1. SQL Serverの配置は、まず必要機能とOSアクセスがRDSの管理境界内で成立するかを確認し、その後ライセンスモデルと運用費を比較する。
2. XはFILESTREAMとホストエージェントのため標準RDSに置けず、確認済みBYOLをEC2 Dedicated Hostsで活用すれば新規SQL Serverライセンス費を避けられる。
3. YはRDS互換でライセンスを所有しないため、License IncludedとMulti-AZを利用し、複数小規模DBの集約でライセンス・基盤・運用の総額を抑える組合せが適切である。

4 正解が要件を満たす理由

Cはワークロードごとに制約を分離する。XはEC2 Dedicated HostsでOSとSQL Serverを完全に制御でき、既存の適格なコアライセンスを活用してFILESTREAMと第三者エージェントを維持する。自己管理コストは残るが、Xを無理にRDSへ置く機能不適合や新規ライセンス費を避けられる。YはRDSのLicense Includedで別途ライセンス購入をせず、Multi-AZ、バックアップ、パッチを管理サービスへ委ね、20DBを適切に集約して運用費も下げられる。

5・6 各選択肢の検討

不適切な理由と、正解になり得る条件を対にして確認します。

A 不正解

RDS for SQL Serverは直接のホストアクセスを提供せず、FILESTREAMもサポートしない。マスターユーザーはOS管理者や完全なsysadminではないため、Xの必須機能を導入できない。

この選択肢が正解になる条件

XがFILESTREAMやOSエージェントを廃止でき、すべての必要機能がRDSでサポートされ、License Includedの運用削減が既存ライセンス価値を上回るなら正解になり得る。

B 不正解

EC2はXの制御要件を満たすが、Xの既存ライセンスを捨ててLicense Includedを買い直し、RDSに適するYまで個別自己管理するため、3年間のライセンス・運用総額が増える。

この選択肢が正解になる条件

全DBがOSレベルのカスタマイズを必要とし、既存ライセンスを利用できず、RDSの機能制限や管理境界を受け入れられないならEC2 License Includedが正解になり得る。

C 正解

Xの機能・既存ライセンス制約にはEC2 BYOL、Yの管理性・ライセンス不在にはRDS License Includedを割り当て、総額を最適化する。

D 不正解

長期的なライセンス削減余地はあるが、SQL、FILESTREAM、エージェントを含む異種エンジン改修が必要で、6か月以内にアプリケーション変更できないという移行制約に反する。

この選択肢が正解になる条件

十分な改修期間と試験予算があり、SQL Server固有機能を廃止してオープンソース互換エンジンへ移行する長期TCO削減が優先なら正解になり得る。

7 今後使える判断基準

商用DBのコスト問題では、時間単価の前に必須機能、OS・sysadminアクセス、ライセンス持込み条件を確認する。管理サービスに適合するDBはRDSで運用費を下げ、非対応機能や有効なBYOL価値があるDBだけをEC2へ残す分割配置も検討する。

8 関連サービスとの比較

RDS for SQL ServerのLicense IncludedはSQL Serverライセンス、基盤、管理機能を料金に含み、Multi-AZやバックアップ運用を軽減するが、ホストアクセスと一部機能に制限がある。EC2は完全な制御と適格なBYOLを選ぶ一方、HA、パッチ、バックアップを自ら管理する。Auroraへの異種移行はライセンスを減らせるが、変換・試験・改修費もTCOに含める。

QUESTION 100

問題100

SECURE

本番相当

単一選択

RAMによる共有VPCと集中接続

企業はOrganizations配下に30のアプリケーションアカウントを持つ。ネットワークアカウントは3AZの非デフォルトVPC、各AZのプライベートサブネット、集中型egress、オンプレミス接続済みTransit Gateway（TGW）を管理している。新規アプリは各自のアカウント所有のEC2、RDS、Lambdaをこの共有ネットワークへ配置し、費用・IAM・サービスクォータの分離を維持したい。アプリチームには自分のリソースとセキュリティグループを管理させるが、VPC、サブネット、ルートテーブル、ネットワークACL、egress経路はネットワークチームだけが変更する。一部の既存アプリは自分のVPCを維持し、それを中央TGWへ接続する。アカウント追加時の招待承認とVPCピアリングのメッシュは避けたい。最も適切な構成はどれか。

- A** 全アプリチームへネットワークアカウントのクロスアカウント管理ロールを渡し、そのロールでEC2とRDSをネットワークアカウントに作成する。タグで費用とクォータを論理分離する。
- B** 各アプリケーションアカウントに同一CIDR構成のVPC、NAT Gateway、オンプレミスVPNを作成する。中央VPCとは必要な組合せだけVPCピアリングし、ルートを各チームが管理する。
- C** AWS RAMで中央TGWだけを全アプリケーションアカウントへ共有する。共有サブネットは作らず、新規アプリのリソースはTGWのARNをサブネットIDとして指定して直接起動する。
- D** AWS RAMのOrganizations共有を有効にし、ネットワークアカウントから新規アプリOUへプライベートサブネットを共有する。中央TGWも必要なアカウントへ共有し、既存VPCのアタッチを許可する。

ここで答えを決める

正解・解説は次のページから

ANSWER 100

問題100の正解・解説

1 正解 **D**

- D** AWS RAMのOrganizations共有を有効にし、ネットワークアカウントから新規アプリOUへプライベートサブネットを共有する。中央TGWも必要なアカウントへ共有し、既存VPCのアタッチを許可する。

2 問題文の重要な条件

- アプリリソースの所有アカウントを分けたまま、中央管理VPCのサブネットへ配置したい
- VPC、サブネット、経路、ネットワークACLはネットワークアカウントが管理し、参加者は自身のリソースを管理する
- 既存の個別VPCも共有TGWへ接続し、Organizations内で招待やピアリングメッシュを避けたい

3 正解へ至る判断過程

1. VPC共有ではVPC全体を移管するのではなく、所有者がAWS RAMで非デフォルトサブネットを組織内アカウントまたはOUへ共有する。
2. 参加アカウントは共有サブネットに自分が所有するEC2、RDS、Lambdaなどを作成でき、VPC所有者はサブネット、ルートテーブル、ネットワークACLなどを中央管理できる。
3. 既存の独立VPCは、RAMで共有されたTGWへ各所有アカウントからアタッチすることで、VPCピアリングの組合せ爆発を避け、両方の配置モデルを満たせる。

4 正解が要件を満たす理由

DはAWS RAMの2種類の共有を使い分ける。サブネット共有により、新規アプリのリソースは各アプリケーションアカウントの所有・請求・IAM境界を保ったまま中央VPCへ配置され、ネットワークアカウントはVPC、ルート、NACL、egressを管理できる。TGW共有により、個別VPCを残すアカウントも中央接続へアタッチできる。Organizations共有を有効にすればOUへの共有が自動反映され、組織内の個別招待も不要になる。

5・6 各選択肢の検討

不適切な理由と、正解になり得る条件を対にして確認します。

A 不正解

リソースをネットワークアカウントに作成すると、所有権、請求、サービスクォータが同じアカウントへ集まり、求めるアカウント分離を失う。広いクロスアカウントロールも最小権限に反する。

この選択肢が正解になる条件

単一のプラットフォームアカウントがアプリ資源も含めて所有・運用し、チーム分離をタグとIAMロールだけで行う方針なら正解になり得る。

B 不正解

各アカウントのVPCは強い分離を提供するが、NAT、VPN、ルートの重複費用と運用を増やし、VPCピアリングはアカウント増加時にメッシュ化する。集中egressと中央統制の条件に合わない。

この選択肢が正解になる条件

各チームがネットワーク経路まで完全に独立管理し、重複費用を許容し、接続先が少なくピアリングの非推移性が問題にならないなら正解になり得る。

C 不正解

TGW共有は参加アカウントのVPCを接続できるが、TGW自体はリソースを起動するサブネットではない。新規アプリが中央VPCへ直接配置される要件にはサブネット共有が必要である。

この選択肢が正解になる条件

すべてのアプリが各自のVPCを所有し続け、必要なのがそれらを中央のオンプレミス・egress経路へ接続することだけならTGW共有が正解になり得る。

D 正解

サブネット共有で新規資源の所有権分離とネットワーク集中管理を両立し、TGW共有で既存個別VPCも統合する。

7 今後使える判断基準

複数アカウントで中央VPCへ各アカウント所有の資源を直接置くならRAMのサブネット共有、各アカウント所有VPCを中央ハブへ接続するならRAMのTGW共有を使う。共有対象、資源所有者、経路管理者を区別して選ぶ。

8 関連サービスとの比較

VPCサブネット共有は所有者がVPC、サブネット、ルート、NACLを中央管理し、参加者が共有サブネット内の自分のリソースを所有する。TGW共有は参加者の独立VPCをハブヘアタッチする。VPCピアリングは少数接続では簡潔だが非推移的で大規模メッシュに不向きであり、クロスアカウントロールは資源共有ではなく別アカウントでの操作委任である。

範囲と注意事項

本書はAWS公式のSAA-C03試験ガイドに示された4ドメインと、試験範囲内サービス一覧を基準に作成した独自問題集です。AWS公式問題の転載ではありません。

参照:

[AWS Certified Solutions Architect - Associate Exam Guide \(SAA-C03\)](#)

[In-Scope AWS Services](#)

[AWS Well-Architected Framework](#)

AWSの機能、制限、価格は変更されることがあります。本書は2026年7月時点の設計判断を前提としています。AWS、各サービス名および関連する商標は、それぞれの権利者に帰属します。