



AWS Glue および Amazon EMR Spark ジョブの Spark SQL クエリのチューニング

AWS 規範ガイドンス



AWS 規範ガイド: AWS Glue および Amazon EMR Spark ジョブの Spark SQL クエリのチューニング

Copyright © 2025 Amazon Web Services, Inc. and/or its affiliates. All rights reserved.

Amazon の商標およびトレードドレスは Amazon 以外の製品およびサービスに使用することはできません。また、お客様に誤解を与える可能性がある形式で、または Amazon の信用を損なう形式で使用することもできません。Amazon が所有していないその他のすべての商標は Amazon との提携、関連、支援関係の有無にかかわらず、それら該当する所有者の資産です。

Table of Contents

序章	1
対象者	1
アーキテクチャ	2
ベストプラクティス	6
結合ヒントの使用	6
ブロードキャスト	6
MERGE	8
SHUFFLE_HASH	9
SHUFFLE_REPLICATE_NL	11
SQL ウィンドウ関数の使用	11
キャッシュ時の列形式の使用	12
パーティショニングヒントの使用	12
Adaptive Query Execution を使用する	12
動的パーティションのプルーニング	13
コストベースオブティマイザーの使用	15
列指向形式の使用	16
よくある質問	17
を使用してApache Sparkジョブをスケーリングし、データをパーティション化するためのベ ストプラクティスは何ですか AWS Glue ?	17
Amazon EMR または でSparkアプリケーションを送信する際に発生するout-of-memoryエラー を解決し、これらのエラーを防ぐにはどうすればよいですか AWS Glue ?	17
でメモリと CPU を管理するにはどうすればよいですか AWS Glue ?	17
リソース	18
AWS ドキュメント	18
Spark ドキュメント	18
ドキュメント履歴	19
用語集	20
#	20
A	21
B	24
C	26
D	29
E	33
F	35

G	36
H	38
I	39
L	41
M	42
O	46
P	49
Q	52
R	52
S	55
T	59
U	60
V	61
W	61
Z	62
.....	lxiii

AWS Glue および Amazon EMR Spark ジョブの Spark SQL クエリの調整

Phani Alapaty と Ravikiran Rao、Amazon Web Services (AWS)

2024 年 1 月 ([ドキュメント履歴](#))

[Spark SQL](#) は、構造化データを処理するための Apache Spark モジュールです。[Amazon EMR](#) および [AWS Glue](#) ジョブは Spark SQL、を使用してデータの処理、変換、ロードを行います。基本的な Spark 回復力のある分散データセット (RDD) API とは異なり、Spark SQL インターフェイスはデータと実行される計算の両方の構造 Spark に関する詳細情報を提供します。内部的には、この追加情報 Spark SQL を使用して追加のクエリ最適化を実行します。SQL やデータセット API など Spark SQL、を操作する方法はいくつかあります。

データの結合は、オブジェクトストアまたはデータベースにデータを抽出、変換、またはロードするときに実行できる、最も一般的で重要なオペレーションの 1 つです。参加するときは、パフォーマンスを考慮する必要があります。結合、分析、または集計オペレーションの一部がメモリ不足になる場合、大規模なネットワーク転送など、複数のシナリオがあります。これにより、ジョブが AWS Glue Spark 失敗する可能性があります。

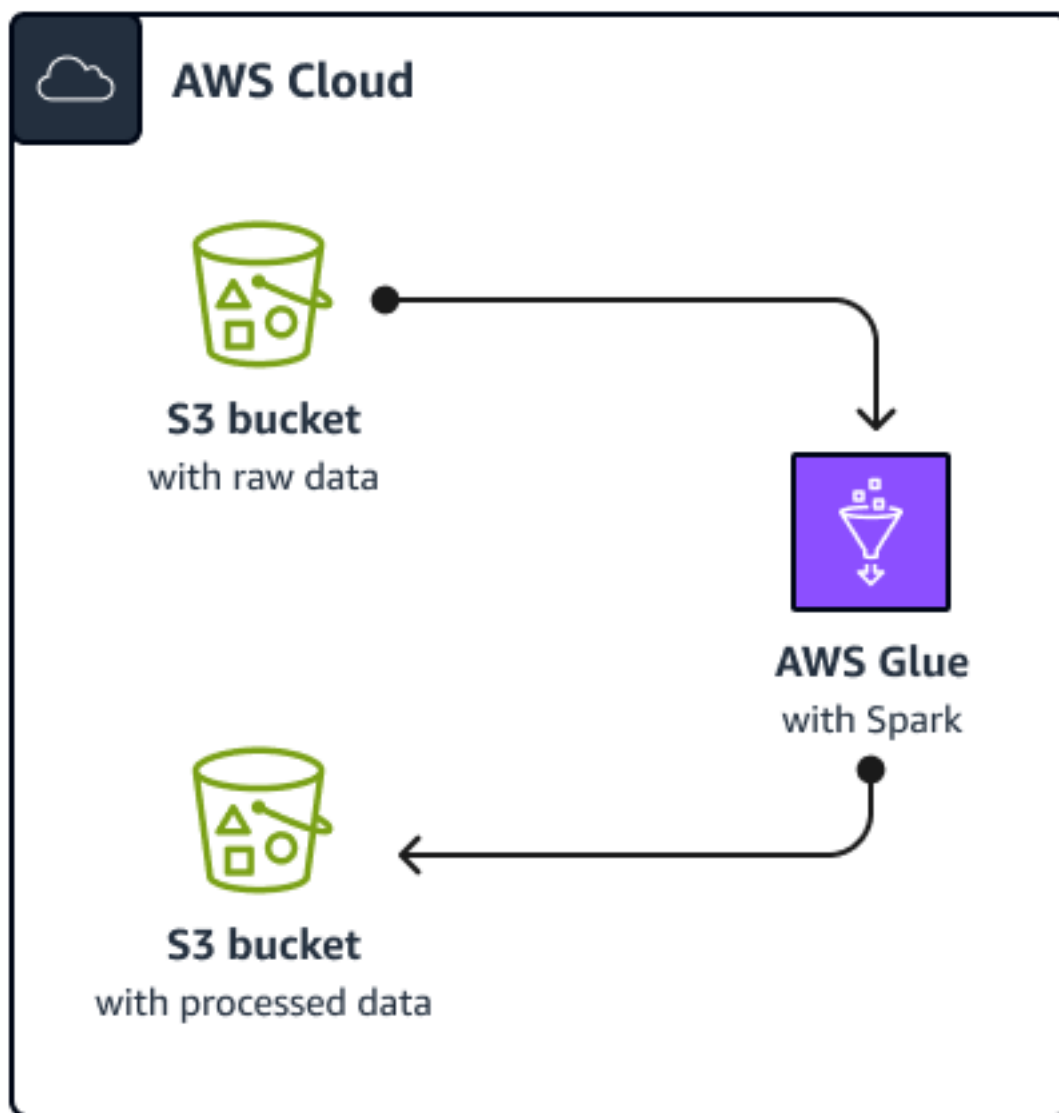
このガイドでは、AWS Glue または Amazon EMR ジョブの Spark SQL 結合クエリを調整するのに役立つベストプラクティスを提供します。Spark には、Spark SQL ワークロードのパフォーマンスを向上させる多くの設定オプションが用意されています。これらの調整はプログラムで行うことも、`spark-submit` コマンドを使用してグローバルレベルで適用することもできます。このガイドでは、Spark SQL クエリとアプリケーションのパフォーマンスを改善または微調整できるように、これらの設定の一部について説明します。このガイドの推奨事項は、Spark SQL クエリとアプリケーションのパフォーマンスを向上させるためにプロフェッショナルサービスが AWS 使用する設定に基づいています。

対象者

このガイドは、アーキテクト、データエンジニア、データサイエンティスト、デベロッパーが、Spark SQL クエリのパフォーマンスを向上させる Spark SQL 設定オプションを理解するのに役立ちます。

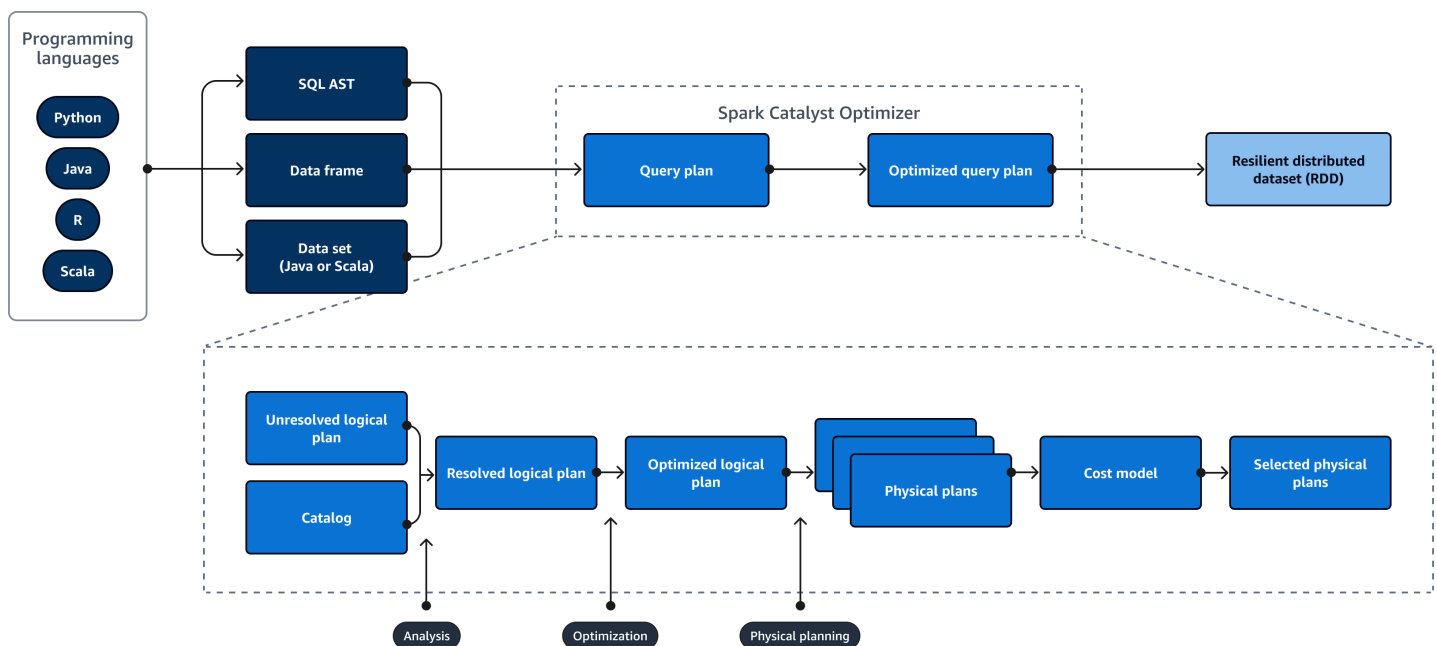
アーキテクチャ

次の図は、このガイドで説明されているソリューションのアーキテクチャを示しています。AWS Glue ジョブは、[Amazon Simple Storage Service \(Amazon S3\)](#) バケットからデータを読み取ります。バケットは、データの保存、保護、取得に役立つクラウドベースのオブジェクトストレージサービスです。ジョブは AWS Glue Spark SQL AWS Management Console、[AWS Command Line Interface \(AWS CLI\)](#)、または AWS Glue API を使用して開始できます。ジョブは、Amazon AWS Glue Spark SQL S3 バケット内の raw データを処理し、処理されたデータを別のバケットに保存します。Amazon S3



例として、このガイドでは、Pythonと Spark SQL () で記述された基本的な[AWS Glue Spark SQL ジョブ](#)について説明しますPySpark。この AWS Glue ジョブは、Spark SQLチューニングのベストプラクティスを示すために使用されます。このガイドでは に焦点を当てていますが AWS Glue、このガイドのベストプラクティスは [Amazon EMR Spark SQLジョブ](#)にも適用されます。

次の図は、Spark SQLクエリのライフサイクルを示しています。Spark SQL Catalyst Optimizer はクエリプランを生成します。クエリプランは、SQL リレーショナルデータベースシステムのデータにアクセスするために使用する手順などの一連のステップです。パフォーマンス最適化Spark SQLクエリプランを開発するには、まずEXPLAINプランを表示し、プランを解釈してから、プランを調整します。Spark SQL ユーザーインターフェイス (UI) またはSpark SQL履歴サーバーを使用して、プランを視覚化できます。



Spark Catalyst Optimizer は、次のように初期クエリプランを最適化されたクエリプランに変換します。

1. 分析と宣言 APIs – 分析フェーズが最初のステップです。SQL クエリで参照されるオブジェクトが不明であるか、入力テーブルと一致していない未解決の論理プランは、バインドされていない属性とデータ型で生成されます。次に、Spark SQL Catalyst Optimizer は一連のルールを適用して論理プランを構築します。SQL パーサーは SQL Abstract Syntax Tree (AST) を生成し、これを論理プランの入力として提供できます。入力、API を使用して構築されたデータフレームまたはデータセットオブジェクトである場合もあります。次の表は、SQL、データフレーム、またはデータセットを使用するタイミングを示しています。

	SQL	データフレーム	データセット
構文エラー	ランタイム	コンパイル時間	コンパイル時間
分析エラー	ランタイム	ランタイム	コンパイル時間

入力のタイプの詳細については、以下を参照してください。

- データセット API は型付きバージョンを提供します。これにより、ユーザー定義の Lambda 関数に大きく依存するため、パフォーマンスが低下します。RDD またはデータセットは静的に型付けされます。たとえば、RDD を定義する場合は、スキーマ定義を明示的に指定する必要があります。
 - データフレーム API は、型のないリレーショナルオペレーションを提供します。データフレームは動的に入力されます。RDD と同様に、データフレームを定義すると、スキーマは同じままになります。データは構造化されたままです。ただし、この情報は実行時にのみ使用できます。これにより、コンパイラは SQL のようなステートメントを記述し、新しい列をその場で定義できます。たとえば、オペレーションごとに新しいクラスを定義することなく、既存のデータフレームに列を追加できます。
 - Spark SQL クエリはランタイム中に構文エラーと分析エラーについて評価されるため、ランタイムが高速化されます。
- カタログ – Spark SQL は、データベース、テーブル、列、パーティションなどの永続的なリレーショナルエンティティのメタデータを管理する Apache Hive Metastore (HMS) ために を使用します。
 - 最適化 — オプティマイザはヒューリスティックとコストを使用してクエリプランを書き換えます。最適化された論理プランを作成するには、以下を実行します。
 - 列のプルーニング
 - 述語をプッシュダウンする
 - 結合の順序を変更する
 - 物理プランとプランナー – Spark SQL Catalyst Optimizer は論理プランを一連の物理プランに変換します。つまり、 の内容を変換します。
 - 選択した物理プラン — Spark SQL Catalyst Optimizer は、最も費用対効果の高い物理プランを選択します。
 - 最適化クエリプラン – パフォーマンス最適化およびコスト最適化クエリプラン Spark SQL を実行します。Spark SQL Memory Management は、メモリ使用量を追跡し、タスクとオペレータ間で

メモリを分散します。Spark SQL Tungsten エンジンでは、Spark SQLアプリケーションのメモリとCPU 効率を大幅に向上させることができます。また、バイナリデータモデル処理を実装し、バイナリデータに対して直接動作します。これにより、逆シリアル化の必要性が回避され、データ変換と逆シリアル化に関連するオーバーヘッドが大幅に削減されます。

ベストプラクティス

このセクションでは、Amazon EMR および の Spark SQL クエリを調整するための以下のベストプラクティスについて説明します AWS Glue。

- [での結合ヒントの使用 Spark SQL](#)
- [join および groupBy の代わりに SQL ウィンドウ関数を使用する](#)
- [キャッシュ時の列形式の使用](#)
- [Spark3.0.0 でのパーティショニングヒントの使用](#)
- [Adaptive Query Execution を使用する](#)
- [動的パーティションのプルーニング](#)
- [Spark コストベースのオプティマイザの使用](#)
- [列形式の使用によるクエリパフォーマンスの向上](#)

での結合ヒントの使用 Spark SQL

Spark3.0 では、実行時に Spark 使用する結合アルゴリズムのタイプを指定できます。結合戦略ヒント、`BROADCASTMERGESHUFFLE_HASH`、および `SHUFFLE_REPLICATE_NL`、指定された各リレーションを別のリレーションに結合するときに、ヒント戦略を使用する Spark ように に指示します。このセクションでは、結合戦略のヒントについて詳しく説明します。

ブロードキャスト

Broadcast Hash 結合では、データセットの 1 つが他のデータセットよりも大幅に小さくなります。小さいデータセットはメモリに収まるため、クラスター内のすべてのエグゼキューターにブロードキャストされます。データがブロードキャストされると、標準のハッシュ結合が実行されます。Broadcast Hash 結合は 2 つのステップで行われます。

- Broadcast – 小さいデータセットは、クラスター内のすべてのエグゼキューターにブロードキャストされます。
- ハッシュ結合 – 小さいデータセットはすべてのエグゼキューターでハッシュされ、大きいデータセットに結合されます。

`sort` または `merge` オペレーションはありません。大きなファクトテーブルを、スタースキーマ結合の実行に使用される小さなディメンションテーブルで結合する場合、Broadcast Hash

は最速の結合アルゴリズムです。次の例は、ブロードキャストハッシュ結合の仕組みを示しています。ヒントを含む結合側は、`spark.sql.autoBroadcastJoinThreshold` プロパティで指定されたサイズ制限に関係なくブロードキャストされます。結合の両側にブロードキャストヒントがある場合、サイズが小さいもの (統計に基づく) がブロードキャストされます。`spark.sql.autoBroadcastJoinThreshold` プロパティのデフォルト値は 10 MB です。これにより、結合の実行時にすべてのワーカーノードにブロードキャストされるテーブルの最大サイズがバイト単位で設定されます。

次の例では、クエリ、物理 EXPLAIN プラン、およびクエリの実行にかかる時間を示します。2 番目の EXPLAIN 計画例に示すように、`BROADCASTJOIN` ヒントを使用してブロードキャスト結合を強制する場合、クエリの処理時間を短縮できます。

```
SQL Query : select table1.id,table1.col,table2.id,table2.int_col from table1 join
table2 on table1.id = table2.id
```

```
== Physical Plan ==
```

```
AdaptiveSparkPlan isFinalPlan=false
```

```
+ - SortMergeJoin [id#80L], [id#95L], Inner
  :- Sort [id#80L ASC NULLS FIRST], false, 0
  :   +- Exchange hashpartitioning(id#80L, 36), ENSURE_REQUIREMENTS, [id=#725]
  :     +- Filter isnotnull(id#80L)
  :       +- Scan ExistingRDD[id#80L,col#81]
+ - Sort [id#95L ASC NULLS FIRST], false, 0
  +- Exchange hashpartitioning(id#95L, 36), ENSURE_REQUIREMENTS, [id=#726]
    +- Filter isnotnull(id#95L)
      +- Scan ExistingRDD[id#95L,int_col#96L]
```

```
Number of records processed: 799541
```

```
Querytime : 21.87715196 seconds
```

```
SQL Query :
```

```
select /*+ BROADCASTJOIN(table1)*/ table1.id,table1.col,table2.id,table2.int_col from
table1 join table2 on table1.id = table2.id
```

```
Physical Plan ==
```

```
AdaptiveSparkPlan isFinalPlan=false\n
```

```
+ - BroadcastHashJoin [id#271L], [id#286L], Inner, BuildLeft, false
  :- BroadcastExchange HashedRelationBroadcastMode(List(input[0, bigint,
false]),false), [id=#955]
  :   +- Filter isnotnull(id#271L)
  :     +- Scan ExistingRDD[id#271L,col#272]
```

```
+-- Filter isnotnull(id#286L)
+-- Scan ExistingRDD[id#286L,int_col#287L]
```

Number of records processed: 799541

Querytime : 15.35717314 seconds

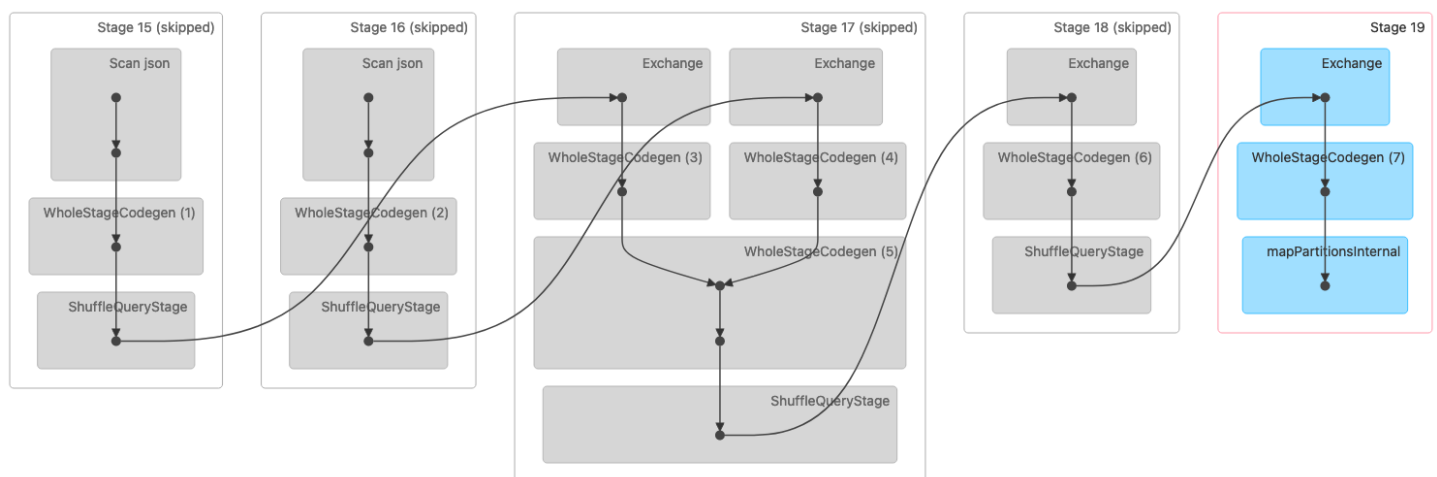
MERGE

両方のデータセットが大きく、メモリに収まらない場合は、シャッフルソートマージ結合が推奨されます。名前が示すように、この結合には次の 3 つのフェーズが含まれます。

- シャッフルフェーズ – 結合クエリの両方のデータセットがシャッフルされます。
- ソートフェーズ – レコードは両側の結合キーでソートされます。
- マージフェーズ – 結合条件の両側は、結合キーに基づいて反復されます。

次の画像は、シャッフルソートマージ結合の有向非巡回グラフ (DAG) の視覚化を示しています。両方のテーブルは最初の 2 つのステージで読み取られます。次のステージ (ステージ 17) では、シャッフルされ、ソートされ、最後にマージされます。

注: このイメージの一部のステージは、これらのステップが前のステージで完了したため、スキップ済みとして表示されます。そのデータはキャッシュまたは保持され、これらのステージで使用されます。



以下は、ソートマージ結合を示す物理プランです。

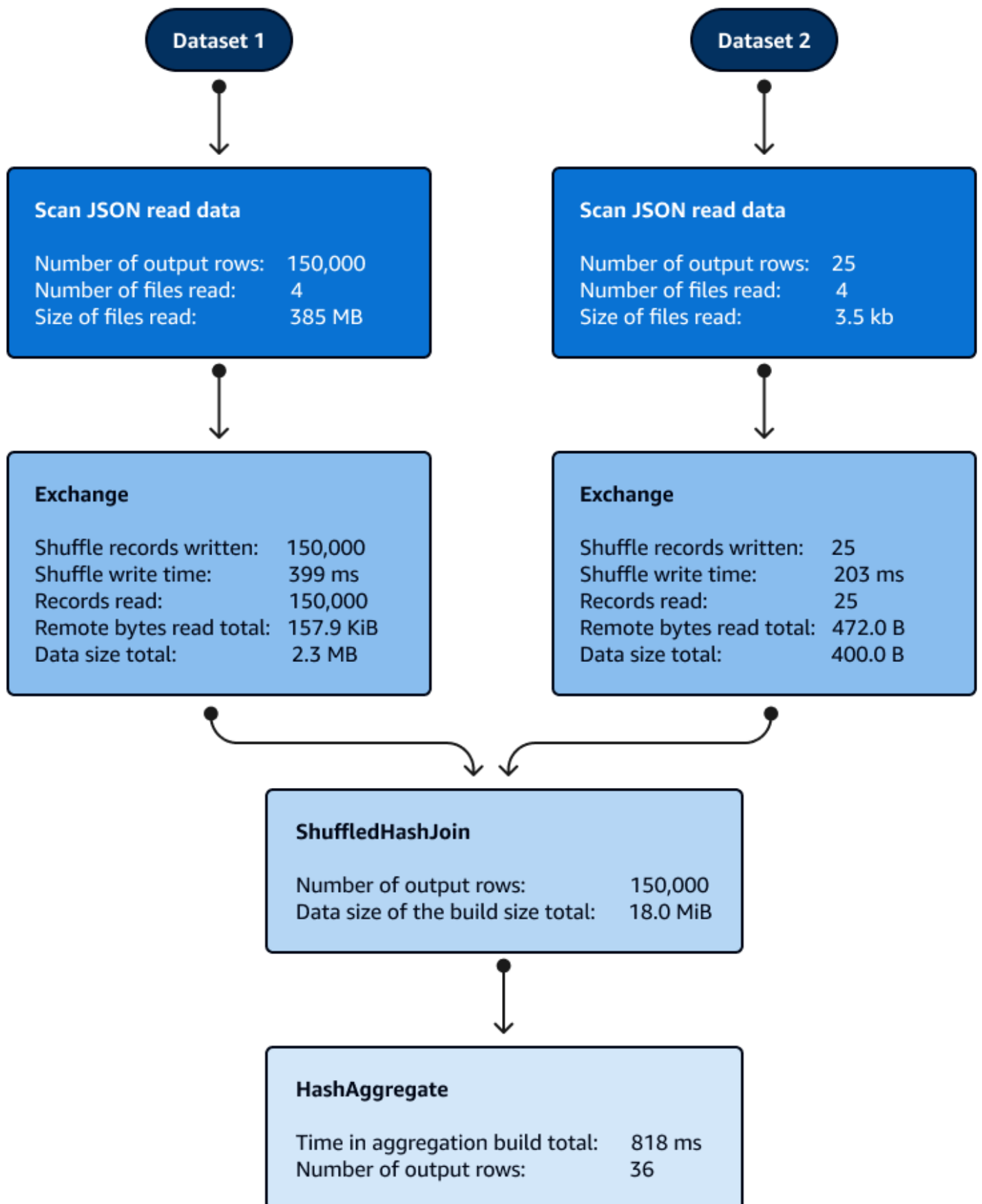
```
== Physical Plan ==
AdaptiveSparkPlan isFinalPlan=false
  +- SortMergeJoin [id#320L], [id#335L], Inner
    :- Sort [id#320L ASC NULLS FIRST], false, 0
    :   +- Exchange hashpartitioning(id#320L, 36), ENSURE_REQUIREMENTS, [id=#1018]
    :     +- Filter isnotnull(id#320L)
    :       +- Scan ExistingRDD[id#320L,col#321]
    :         +- Sort [id#335L ASC NULLS FIRST], false, 0
    :           +- Exchange hashpartitioning(id#335L, 36), ENSURE_REQUIREMENTS,
[id=#1019]
    :             +- Filter isnotnull(id#335L)
    :               +- Scan ExistingRDD[id#335L,int_col#336L]
```

SHUFFLE_HASH

シャッフルハッシュ結合は、名前が示すように、両方のデータセットをシャッフルすることで機能します。両側の同じキーは、同じパーティションまたはタスクになります。データがシャッフルされると、2つのデータセットのうち最小のデータセットがバケットにハッシュされ、パーティション内でハッシュ結合が実行されます。データセット全体がブロードキャストされないため、シャッフルハッシュ結合はブロードキャストハッシュ結合とは異なります。シャッフルハッシュ結合は2つのフェーズに分かれています。

- シャッフルフェーズ – 両方のデータセットがシャッフルされます。
- ハッシュ結合フェーズ – データの小さい側がハッシュされ、バケット化され、すべてのパーティションの大きい側とハッシュ結合されます。

パーティション内のシャッフルハッシュ結合ではソートは必要ありません。次の図は、シャッフルハッシュ結合のフェーズを示しています。データは最初に読み取られ、次にシャッフルされ、ハッシュが作成されて結合に使用されます。



デフォルトでは、オプティマイザはブロードキャストハッシュ結合を使用できないときにシャッフルハッシュ結合を選択します。ブロードキャストハッシュ結合のしきい値サイズ (`spark.sql.autoBroadcastJoinThreshold`) と選択したシャッフルパーティションの数 (`spark.sql.shuffle.partitions`) に基づいて、論理 SQL の 1 つのパーティションがローカルハッシュテーブルを構築するのに十分な大きさである場合、シャッフルハッシュ結合を使用します。

SHUFFLE_REPLICATE_NL

デカルト積結合とも呼ばれる Shuffle-and-Replicate ネストされたループ結合は、データセットがブロードキャストされていない点を除いて、ブロードキャストハッシュ結合と非常によく似ています。

この結合アルゴリズムでは、同じキーを持つレコードが同じパーティションに送信されないため、シャッフルは真のシャッフルを参照しません。代わりに、両方のデータセットのパーティション全体がネットワーク経由でコピーされます。データセットのパーティションが使用可能になると、ネストされたループ結合が実行されます。最初のデータセットにレコード X の数があり、各パーティションの 2 番目のデータセットにレコード Y の数がある場合、2 番目のデータセットの各レコードは、最初のデータセットのすべてのレコードに結合されます。これは、すべてのパーティションでループ $X \times Y$ 時間で継続されます。

join および groupBy の代わりに SQL ウィンドウ関数を使用する

[ウィンドウ関数](#) (Spark ドキュメント) は、現在のレコードに関連するウィンドウと呼ばれる行のグループに対して計算を実行します。たとえば、現在の行と同じパーティションまたはフレーム内に配置できます。ウィンドウ関数は Spark に組み込まれており、集計関数と似ています。SUM やなどの集計関数は、行のグループで動作し MAX、グループごとに 1 つの戻り値を計算します。ウィンドウ関数は、移動平均の計算や、現在の行の相対位置に基づく行の値へのアクセスなど、タスクの処理に役立ちます。

Spark groupBy 関数はデータをグループに収集し、グループ化されたデータに対して集計関数を実行します。を使用するとgroupBy、は最初にデータをSpark部分的に集約し、次に縮小されたデータセットをシャッフルします。ウィンドウ関数を使用すると、データセット全体がシャッフルされます。ウィンドウ関数を使用すると、ランタイムを高速化できます。

非常に大きなデータセットでは、列の基数が大きい場合、ウィンドウ関数が推奨されます。ただし、列の基数が少ない場合、データ集約は小さく、集計結果は結合でブロードキャストできます。

キャッシュ時の列形式の使用

Spark SQL には、インメモリ形式のテーブルを列指向形式でキャッシュする機能があります。spark.catalog.cacheTable("tableName")またはdataFrame.cache()関数呼び出しを使用して、インメモリ列指向形式でテーブルをキャッシュできます。その後、Spark SQLエンジンは必要な列のみをスキャンし、圧縮を自動的に調整してメモリと CPU 使用率を減らします。spark.catalog.uncacheTable("tableName")またはを使用してdataFrame.unpersist()、メモリからテーブルを削除できます。

Spark3.0.0 でのパーティショニングヒントの使用

Spark [パーティショニングヒント](#)は、パフォーマンスを調整し、出力ファイルの数を減らすのに役立ちます。は、COALESCE、REPARTITIONなどのパーティショニングヒントSpark SQLをサポートしていますREPARTITION_BY_RANGE。これらのヒントは、repartitionなどのデータセット APIs coalesceに似ていますrepartitionByRange。以下のヒントは、で出力ファイルの数を制御するのに役立ちます。これによりSpark SQL、パフォーマンスを調整できます。

- Coalesce - パーティションの数を指定されたパーティションの数に減らします。パーティション番号は、COALESCEヒントの唯一のパラメータです。
- 再パーティション - 指定されたパーティション式を使用して、指定された数のパーティションに再パーティションします。REPARTITION ヒントパラメータは、パーティション番号、列名、またはその両方です。
- 範囲による再パーティション - 指定されたパーティション式を使用して、指定された数のパーティションに再パーティションします。列名はREPARTITION_BY_RANGEヒントの必須パラメータであり、パーティション番号はオプションです。
- 再調整 - すべてのパーティションが適切なサイズになるように、クエリ結果の出力パーティションを再調整します。REBALANCEヒントパラメータは、最初のパーティション番号、列名、またはその両方です。

Adaptive Query Execution を使用する

Spark3.0 以降には、と呼ばれる追加の最適化レイヤーが含まれています[Adaptive Query Execution \(AQE\)](#)。クエリの実行中に収集されたメトリクスに基づいてクエリを最適化します。ランタイム統計を使用して、最も効率的な実行計画を選択します。この機能は、Apache Sparkバージョン3.2.0 以降ではデフォルトで有効になっています。Spark 設定を に設定することでオフにできますspark.sql.adaptive.enabled=false。AQE には、シャッフル後のパーティションの結合、

ソートマージ結合のブロードキャスト結合への変換、スキュー結合の最適化の 3 つの主な機能があります。

動的パーティションのプルーニング

Spark3.0 以降には、動的パーティションプルーニング (DPP) が含まれています。動的パーティションプルーニングは、データの読み取り時に不要なパーティションのスキューを防ぐ Spark の最適化手法です。以下は、DPP について知っておくべき重要な点です。

- クエリフィルターと述語でリクエストされたパーティション値を調べ、クエリを満たすために必要なパーティションを決定します。不要と見なされるパーティションは、自動的に透過的にプルーニングされます。
- DPP は、該当するデータを含まないパーティションをスキップすることで、処理時間とリソース使用率を削減します。これにより Spark、関連するパーティションのみに集中できます。
- 挿入または増分ロードによって追加される静的パーティションと動的に生成されたパーティションの両方で動作します。新しいパーティション Spark を認識し、動的プルーニングの適用を継続できます。
- DPP は完全に透過的であるか、開発者には見えません。DPP を有効にするために特別なコーディングは必要ありません。これは、クエリプランの生成中に最適化としてバックグラウンドで自動的に行われます。

以下は、DPP が効率的に機能するようにするためのベストプラクティスです。

- Spark データフレームオペレーションの早い段階でフィルターを適用して、述語のプッシュダウンを使用します。これにより Spark、パーティションメタデータを使用してパーティションを早期に排除できます。
- を `ANALYZE TABLE` 頻繁に実行して、データの統計を収集します。これにより、無視できるパーティションをより正確に判断 Spark するのに役立つ列レベルの統計が削減されます。
- データを過剰に分割しないようにします。パーティションが多すぎると、統計を収集するときにドライバードが過負荷になる可能性があります。大きなテーブルごとに 10 ~ 100 個のパーティションを目指します。
- 結合前にデータフレームを再パーティション化します。これにより、すべてのデータの移動を必要とするシャッフル結合を防ぎ、読み取りデータの量をさらに最適化できます。
- 結合されているさまざまなテーブルで、一貫したパーティション列タイプと命名を使用します。これにより、結合の最適化のためにパーティションをより Spark 適切に一致させることができます。

- クエリをテストEXPLAINして DPP が適用されていることを確認し、追加チューニングが必要かどうかを確認します。

スタースキーマでは、テーブルはファクトテーブルとディメンションテーブルの 2 つの主要なタイプに分割されます。ディメンションテーブルは、ファクトテーブルよりもはるかに小さい傾向があります。ファクトテーブルをディメンションテーブルに結合すると、DPP はクエリプランを最適化します。ディメンションテーブルに適用されるフィルターからサブクエリを作成します。このサブクエリをブロードキャストし、そこからハッシュテーブルを構築します。次に、ファクトテーブルのデータを読み取る前に、ハッシュテーブルをファクトテーブルのスキャンフェーズに適用します。これにより、DPP はより大きなファクトテーブルから読み取る必要があるデータの量を減らすことができます。

次のクエリ例は、DPP が動作していることを示しています。クエリは国 (インド) からの注文数を取得し、ファクトテーブル (fact_orders) とディメンションテーブル (nation) 間の内部結合を含め、nation。fact_orders テーブルは列によってパーティション化されます o_nationkey。

```
- "select n.n_name as country, count(1) as no_of_orders from fact_orders o join nation
  n on o.o_nationkey = n.n_nationkey where n.n_name = 'INDIA' group by n.n_name"
```

以下は、EXPLAINプランで使用されるステップです。

1. 小さいディメンションテーブル (nation) をスキャンし、列でフィルタリングします n_name = 'INDIA'。
2. 前のステップの結果をブロードキャストします。
3. 最初のステップの結果をフィルタリングするサブクエリを作成します。
4. フルテーブルスキャンではなく、必要なファクトテーブルパーティションのみをスキャン PartitionFilterするように、としてプッシュダウンします。

この DPP 最適化クエリの EXPLAIN 計画を次に示します。

```
== Physical Plan ==
AdaptiveSparkPlan isFinalPlan=true
+- == Final Plan ==
   *(4) HashAggregate(keys=[], functions=[count(1)], output=[count#208L])
  +- ShuffleQueryStage 3
     +- Exchange SinglePartition, ENSURE_REQUIREMENTS, [id=#353]
        +- *(3) HashAggregate(keys=[], functions=[partial_count(1)],
          output=[count#212L])
```

```

+- *(3) HashAggregate(keys=[n_name#31], functions=[], output=[])
+- ShuffleQueryStage 1
+- Exchange hashpartitioning(n_name#31, 36), ENSURE_REQUIREMENTS,
[id=#315]
+- *(2) HashAggregate(keys=[n_name#31], functions=[],
output=[n_name#31])
+- *(2) Project [n_name#31]
+- *(2) BroadcastHashJoin [cast(o_nationkey#145 as bigint)],
[n_nationkey#32L], Inner, BuildRight, false
:- *(2) ColumnarToRow
: +- FileScan parquet [o_nationkey#145] Batched:
true, DataFilters: [], Format: Parquet, Location: InMemoryFileIndex[s3://
aws-spark-tuning/fact_orders], PartitionFilters: [isnotnull(o_nationkey#145),
dynamicpruningexpression(cast(o_nationkey#145 as bigint) IN dynamiccp...,
PushedFilters: [], ReadSchema: struct<>
: +- SubqueryBroadcast dynamicpruning#210, 0,
[n_nationkey#32L], [id=#200]
: +- OutputAdapter [n_name#31, n_nationkey#32L]
: +- AdaptiveSparkPlan isFinalPlan=true
: +- BroadcastQueryStage 2
: +- ReusedExchange [n_name#31,
n_nationkey#32L], BroadcastExchange HashedRelationBroadcastMode(List(input[1, bigint,
false]),false), [id=#233]
+- BroadcastQueryStage 0
+- BroadcastExchange
HashedRelationBroadcastMode(List(input[1, bigint, false]),false), [id=#233]
+- *(1) Filter ((isnotnull(n_name#31) AND
(n_name#31 = INDIA)) AND isnotnull(n_nationkey#32L))
+- FileScan json [n_name#31,n_nationkey#32L]
Batched: false, DataFilters: [isnotnull(n_name#31), (n_name#31 = INDIA),
isnotnull(n_nationkey#32L)], Format: JSON, Location: InMemoryFileIndex[s3://
aws-spark-tuning/input/demo/json/nation], PartitionFilters: [], PushedFilters:
[IsNotNull(n_name), EqualTo(n_name,INDIA), IsNotNull(n_nationkey)], ReadSchema:
struct<n_name:string,n_nationkey:bigint>

```

`o_nationkey` 列にダイレクトフィルターが追加されていなくても、DPP 機能のため、はテーブル全体ではなく、必要なパーティションのみ Spark を自動的にスキャンします。

Spark コストベースのオプティマイザの使用

Spark バージョン 2.09 以降では、コストベースのオプティマイザ (CBO) を使用できます。CBO は、さまざまなテーブル統計に基づいて、クエリの最も安価な実行プランを選択します。CBO は

CPU 使用率と I/O に関してクエリの実行を最適化しようとするため、できるだけ早く返します。ただし、ほとんどの場合、特に統計収集がデータ処理よりもさらに高価である場合、データ統計は一般的にありません。統計が利用可能であっても、古くなっている可能性があります。CBO を使用する際のベストプラクティスを次に示します。

- 結合、フィルター、グループバイなどのオペレーションで使われるすべての列の統計を収集します。これにより、オプティマイザはより良い意思決定を行うことができます。を使用して統計を収集 `ANALYZE TABLE` します。
- データが急速に変化している場合は、統計収集の頻度を増やします。古い統計では、計画が最適化されない可能性があります。
- デフォルトの CBO パラメータから開始し、必要に応じて調整します。積極的な最適化により、ランタイムが増加する可能性があります。
- 最適化計画を強制してパフォーマンスを比較することで、CBO が最適な計画を選択しているかどうかを検証します。オプティマイザのログ記録を有効にして、プランを比較します。
- 最適化を向上させるには、パーティショニング、バケット化、データスキューのヒントと組み合わせて CBO を使用します。パーティション分割は CBO のスケーリングに役立ちます。
- 計画の回帰を早期に検出するには、CBO を有効にした後、クエリのパフォーマンスを継続的にモニタリングします。

列形式の使用によるクエリパフォーマンスの向上

Spark は、`CSV` など `Apache Parquet` `Optimized Row Columnar (ORC)`、さまざまな入力ファイル形式を使用できます。ただし、`は` 内で最適に `Parquet` 動作します `Spark SQL`。ランタイムの高速化、スキャンスループットの向上、ディスク I/O の削減、運用コストの削減を実現します。`は`、最小/最大統計などのプッシュダウンフィルターで `Parquet` ファイル統計データを使用して、役に立たないデータを自動的にフィルタリング `Spark` できます。一方、`parquet Spark` ベクトル化されたリーダーを有効にして `Parquet` ファイルをバッチで読み取ることができます。`Spark SQL` を使用してデータを処理している場合は、可能であれば `Parquet` ファイル形式を使用することをお勧めします。

よくある質問

を使用して Apache Spark ジョブをスケーリングし、データをパーティション化するためのベストプラクティスは何ですか AWS Glue ?

詳細については、AWS ビッグデータブログの「[で Apache Spark ジョブをスケーリングし、データをパーティション化するためのベストプラクティス AWS Glue](#)」を参照してください。

Amazon EMR または で Spark アプリケーションを送信する際に発生する out-of-memory エラーを解決し、これらのエラーを防ぐにはどうすればよいですか AWS Glue ?

詳細については、の [EMR ベストプラクティスガイド](#) を参照してください GitHub。

でメモリと CPU を管理するにはどうすればよいですか AWS Glue ?

詳細については、AWS ビッグデータブログの「[でのメモリ管理の最適化 AWS Glue](#)」を参照してください。

リソース

AWS ドキュメント

- [Apache Sparkジョブのパフォーマンス調整 AWS Glue のベストプラクティス](#)
- [AWS Command Line Interface \(AWS CLI\) サービスドキュメント](#)
- [AWS Glue サービスドキュメント](#)
- [Amazon Simple Storage Service \(Amazon S3\) サービスドキュメント](#)
- [Amazon EMR サービスドキュメント](#)

Spark ドキュメント

- [パフォーマンスチューニング](#)
- [SQL クエリの結合戦略ヒント](#)

ドキュメント履歴

以下の表は、本ガイドの重要な変更点について説明したものです。今後の更新に関する通知を受け取る場合は、[RSS フィード](#) をサブスクライブできます。

変更	説明	日付
初版発行	—	2024 年 1 月 4 日

AWS 規範ガイドの用語集

以下は、AWS 規範ガイドによって提供される戦略、ガイド、パターンで一般的に使用される用語です。エントリを提案するには、用語集の最後のフィードバックの提供リンクを使用します。

数字

7 Rs

アプリケーションをクラウドに移行するための 7 つの一般的な移行戦略。これらの戦略は、ガートナーが 2011 年に特定した 5 Rs に基づいて構築され、以下で構成されています。

- リファクタリング/アーキテクチャの再設計 — クラウドネイティブ特徴を最大限に活用して、俊敏性、パフォーマンス、スケーラビリティを向上させ、アプリケーションを移動させ、アーキテクチャを変更します。これには、通常、オペレーティングシステムとデータベースの移植が含まれます。例: オンプレミスの Oracle データベースを Amazon Aurora PostgreSQL 互換エディションに移行します。
- リプラットフォーム (リフトアンドリシェイプ) — アプリケーションをクラウドに移行し、クラウド機能を活用するための最適化レベルを導入します。例: オンプレミスの Oracle データベースを Oracle 用 Amazon Relational Database Service (Amazon RDS) に移行します AWS クラウド。
- 再購入 (ドロップアンドショップ) — 通常、従来のライセンスから SaaS モデルに移行して、別の製品に切り替えます。例: カスタマーリレーションシップ管理 (CRM) システムを Salesforce.com に移行します。
- リホスト (リフトアンドシフト) — クラウド機能を活用するための変更を加えずに、アプリケーションをクラウドに移行します。例: オンプレミスの Oracle データベースをの EC2 インスタンス上の Oracle に移行します AWS クラウド。
- 再配置 (ハイパーバイザーレベルのリフトアンドシフト) — 新しいハードウェアを購入したり、アプリケーションを書き換えたり、既存の運用を変更したりすることなく、インフラストラクチャをクラウドに移行できます。オンプレミスプラットフォームから同じプラットフォームのクラウドサービスにサーバーを移行します。例: Microsoft Hyper-Vアプリケーションをに移行します AWS。
- 保持 (再アクセス) — アプリケーションをお客様のソース環境で保持します。これには、主要なリファクタリングを必要とするアプリケーションや、お客様がその作業を後日まで延期したいアプリケーション、およびそれらを行き移るためのビジネス上の正当性がないため、お客様が保持するレガシーアプリケーションなどがあります。

- 使用停止 — お客様のソース環境で不要になったアプリケーションを停止または削除します。

A

ABAC

[「属性ベースのアクセスコントロール」](#)を参照してください。

抽象化されたサービス

[「マネージドサービス」](#)を参照してください。

ACID

[「アトミック性」、「整合性」、「分離」、「耐久性」](#)を参照してください

アクティブ - アクティブ移行

(双方向レプリケーションツールまたは二重書き込み操作を使用して) ソースデータベースとターゲットデータベースを同期させ、移行中に両方のデータベースが接続アプリケーションからのトランザクションを処理するデータベース移行方法。この方法では、1 回限りのカットオーバーの必要がなく、管理された小規模なバッチで移行できます。より柔軟ですが、[アクティブ/パッシブ移行](#)よりも多くの作業が必要です。

アクティブ - パッシブ移行

ソースデータベースとターゲットデータベースを同期させながら、データがターゲットデータベースにレプリケートされている間、接続しているアプリケーションからのトランザクションをソースデータベースのみで処理するデータベース移行の方法。移行中、ターゲットデータベースはトランザクションを受け付けません。

集計関数

行のグループで動作し、グループの単一の戻り値を計算する SQL 関数。集計関数の例としては、SUMや などがありますMAX。

AI

[「人工知能」](#)を参照してください。

AIOps

[「人工知能オペレーション」](#)を参照してください。

匿名化

データセット内の個人情報を完全に削除するプロセス。匿名化は個人のプライバシー保護に役立ちます。匿名化されたデータは、もはや個人データとは見なされません。

アンチパターン

繰り返し起こる問題に対して頻繁に用いられる解決策で、その解決策が逆効果であったり、効果がなかったり、代替案よりも効果が低かったりするもの。

アプリケーションコントロール

マルウェアからシステムを保護するために、承認されたアプリケーションのみを使用できるようにするセキュリティアプローチ。

アプリケーションポートフォリオ

アプリケーションの構築と維持にかかるコスト、およびそのビジネス価値を含む、組織が使用する各アプリケーションに関する詳細情報の集まり。この情報は、[ポートフォリオの検出と分析プロセス](#) の需要要素であり、移行、モダナイズ、最適化するアプリケーションを特定し、優先順位を付けるのに役立ちます。

人工知能 (AI)

コンピューティングテクノロジーを使用し、学習、問題の解決、パターンの認識など、通常は人間に関連づけられる認知機能の実行に特化したコンピュータサイエンスの分野。詳細については、「[人工知能 \(AI\) とは何ですか?](#)」を参照してください。

AI オペレーション (AIOps)

機械学習技術を使用して運用上の問題を解決し、運用上のインシデントと人の介入を減らし、サービス品質を向上させるプロセス。AWS 移行戦略での AIOps の使用方法については、[オペレーション統合ガイド](#) を参照してください。

非対称暗号化

暗号化用のパブリックキーと復号用のプライベートキーから成る 1 組のキーを使用した、暗号化のアルゴリズム。パブリックキーは復号には使用されないため共有しても問題ありませんが、プライベートキーの利用は厳しく制限する必要があります。

原子性、一貫性、分離性、耐久性 (ACID)

エラー、停電、その他の問題が発生した場合でも、データベースのデータ有効性と運用上の信頼性を保証する一連のソフトウェアプロパティ。

属性ベースのアクセス制御 (ABAC)

部署、役職、チーム名など、ユーザーの属性に基づいてアクセス許可をきめ細かく設定する方法。詳細については、AWS Identity and Access Management (IAM) ドキュメントの「[ABAC AWS](#)」を参照してください。

信頼できるデータソース

最も信頼性のある情報源とされるデータのプライマリバージョンを保存する場所。匿名化、編集、仮名化など、データを処理または変更する目的で、信頼できるデータソースから他の場所にデータをコピーすることができます。

アベイラビリティーゾーン

他のアベイラビリティーゾーンの障害から AWS リージョン 隔離され、同じリージョン内の他のアベイラビリティーゾーンへの低コストで低レイテンシーのネットワーク接続を提供する 内の別の場所。

AWS クラウド導入フレームワーク (AWS CAF)

のガイドラインとベストプラクティスのフレームワークは、組織がクラウドに成功するための効率的で効果的な計画を立て AWS るのに役立ちます。AWS CAF は、ビジネス、人材、ガバナンス、プラットフォーム、セキュリティ、運用という 6 つの重点分野にガイダンスを整理します。ビジネス、人材、ガバナンスの観点では、ビジネススキルとプロセスに重点を置き、プラットフォーム、セキュリティ、オペレーションの視点は技術的なスキルとプロセスに焦点を当てています。例えば、人材の観点では、人事 (HR)、人材派遣機能、および人材管理を扱うステークホルダーを対象としています。この観点から、AWS CAF は、クラウド導入を成功させるための組織の準備に役立つ人材開発、トレーニング、コミュニケーションに関するガイダンスを提供します。詳細については、[AWS CAF ウェブサイト](#) と [AWS CAF のホワイトペーパー](#) を参照してください。

AWS ワークロード認定フレームワーク (AWS WQF)

データベース移行ワークロードを評価し、移行戦略を推奨し、作業見積もりを提供するツール。AWS WQF は AWS Schema Conversion Tool (AWS SCT) に含まれています。データベーススキーマとコードオブジェクト、アプリケーションコード、依存関係、およびパフォーマンス特性を分析し、評価レポートを提供します。

B

不正なボット

個人または組織に損害を与えることを目的とした[ボット](#)。

BCP

[「事業継続計画」](#)を参照してください。

動作グラフ

リソースの動作とインタラクションを経時的に示した、一元的なインタラクティブビュー。Amazon Detective の動作グラフを使用すると、失敗したログオンの試行、不審な API 呼び出し、その他同様のアクションを調べることができます。詳細については、Detective ドキュメントの[Data in a behavior graph](#)を参照してください。

ビッグエンディアンシステム

最上位バイトを最初に格納するシステム。[エンディアン性](#)も参照してください。

二項分類

バイナリ結果 (2 つの可能なクラスの中の 1 つ) を予測するプロセス。例えば、お客様の機械学習モデルで「この E メールはスパムですか、それともスパムではありませんか」などの問題を予測する必要があるかもしれません。または「この製品は書籍ですか、車ですか」などの問題を予測する必要があるかもしれません。

ブルームフィルター

要素がセットのメンバーであるかどうかをテストするために使用される、確率的でメモリ効率の高いデータ構造。

ブルー/グリーンデプロイ

2 つの異なる同一の環境を作成するデプロイ戦略。現在のアプリケーションバージョンを 1 つの環境 (青) で実行し、新しいアプリケーションバージョンを別の環境 (緑) で実行します。この戦略は、最小限の影響で迅速にロールバックするのに役立ちます。

ボット

インターネット経由で自動タスクを実行し、人間のアクティビティややり取りをシミュレートするソフトウェアアプリケーション。インターネット上の情報のインデックスを作成するウェブクローラーなど、一部のボットは有用または有益です。悪質なボットと呼ばれる他のボットの中には、個人や組織を混乱させたり、損害を与えたりすることを意図したものもあります。

ボットネット

[マルウェア](#)に感染し、[ボット](#)ハーダーまたはボットオペレーターとして知られる単一関係者の管理下にあるボットのネットワーク。ボットは、ボットとその影響をスケールするための最もよく知られているメカニズムです。

ブランチ

コードリポジトリに含まれる領域。リポジトリに最初に作成するブランチは、メインブランチといます。既存のブランチから新しいブランチを作成し、その新しいブランチで機能を開発したり、バグを修正したりできます。機能を構築するために作成するブランチは、通常、機能ブランチと呼ばれます。機能をリリースする準備ができたなら、機能ブランチをメインブランチに統合します。詳細については、「[ブランチの概要](#)」(GitHub ドキュメント)を参照してください。

ブレイクグラスアクセス

例外的な状況では、承認されたプロセスを通じて、ユーザーが AWS アカウント 通常アクセス許可を持たない にすばやくアクセスできるようになります。詳細については、Well-Architected [ガイド](#)の「[ブレイクグラス手順の実装](#)」インジケータ AWS を参照してください。

ブラウンフィールド戦略

環境の既存インフラストラクチャ。システムアーキテクチャにブラウンフィールド戦略を導入する場合、現在のシステムとインフラストラクチャの制約に基づいてアーキテクチャを設計します。既存のインフラストラクチャを拡張している場合は、ブラウンフィールド戦略と[グリーンフィールド](#)戦略を融合させることもできます。

バッファキャッシュ

アクセス頻度が最も高いデータが保存されるメモリ領域。

ビジネス能力

価値を生み出すためにビジネスが行うこと (営業、カスタマーサービス、マーケティングなど)。マイクロサービスのアーキテクチャと開発の決定は、ビジネス能力によって推進できます。詳細については、ホワイトペーパー [AWSでのコンテナ化されたマイクロサービスの実行](#) の [ビジネス機能を中心に組織化](#) セクションを参照してください。

ビジネス継続性計画 (BCP)

大規模移行など、中断を伴うイベントが運用に与える潜在的な影響に対処し、ビジネスを迅速に再開できるようにする計画。

C

CAF

[AWS 「クラウド導入フレームワーク」](#) を参照してください。

Canary デプロイ

エンドユーザーへのバージョンのスローリリースと増分リリース。確信が持てば、新しいバージョンをデプロイし、現在のバージョン全体を置き換えます。

CCoE

[「Cloud Center of Excellence」](#) を参照してください。

CDC

[「データキャプチャの変更」](#) を参照してください。

変更データキャプチャ (CDC)

データソース (データベーステーブルなど) の変更を追跡し、その変更に関するメタデータを記録するプロセス。CDC は、ターゲットシステムでの変更を監査またはレプリケートして同期を維持するなど、さまざまな目的に使用できます。

カオスエンジニアリング

障害や破壊的なイベントを意図的に導入して、システムの耐障害性をテストします。[AWS Fault Injection Service \(AWS FIS \)](#) を使用して、AWS ワークロードにストレスを与え、その応答を評価する実験を実行できます。

CI/CD

[継続的インテグレーションと継続的デリバリー](#) を参照してください。

分類

予測を生成するのに役立つ分類プロセス。分類問題の機械学習モデルは、離散値を予測します。離散値は、常に互いに区別されます。例えば、モデルがイメージ内に車があるかどうかを評価する必要がある場合があります。

クライアント側の暗号化

ターゲットがデータ AWS のサービスを受信する前のローカルでのデータの暗号化。

Cloud Center of Excellence (CCoE)

クラウドのベストプラクティスの作成、リソースの移動、移行のタイムラインの確立、大規模変革を通じて組織をリードするなど、組織全体のクラウド導入の取り組みを推進する学際的なチーム。詳細については、AWS クラウド エンタープライズ戦略ブログの [CCoE 投稿](#) を参照してください。

クラウドコンピューティング

リモートデータストレージと IoT デバイス管理に通常使用されるクラウドテクノロジー。クラウドコンピューティングは、一般的に [エッジコンピューティング](#) テクノロジーに接続されています。

クラウド運用モデル

IT 組織において、1 つ以上のクラウド環境を構築、成熟、最適化するために使用される運用モデル。詳細については、[「クラウド運用モデルの構築」](#) を参照してください。

導入のクラウドステージ

組織が に移行するときに通常実行する 4 つのフェーズ AWS クラウド :

- プロジェクト — 概念実証と学習を目的として、クラウド関連のプロジェクトをいくつか実行する
- 基礎固め — お客様のクラウドの導入を拡大するための基礎的な投資 (ランディングゾーンの作成、CCoE の定義、運用モデルの確立など)
- 移行 — 個々のアプリケーションの移行
- 再発明 — 製品とサービスの最適化、クラウドでのイノベーション

これらのステージは、AWS クラウド エンタープライズ戦略ブログのブログ記事 [「クラウドファーストへのジャーニー」](#) と [「導入のステージ」](#) で Stephen Orban によって定義されました。AWS 移行戦略との関連性については、[「移行準備ガイド」](#) を参照してください。

CMDB

[「設定管理データベース」](#) を参照してください。

コードリポジトリ

ソースコードやその他の資産 (ドキュメント、サンプル、スクリプトなど) が保存され、バージョン管理プロセスを通じて更新される場所。一般的なクラウドリポジトリには、GitHub または が含まれます Bitbucket Cloud。コードの各バージョンはブランチと呼ばれます。マイクロサービスの構造では、各リポジトリは 1 つの機能専用です。1 つの CI/CD パイプラインで複数のリポジトリを使用できます。

コールドキャッシュ

空である、または、かなり空きがある、もしくは、古いデータや無関係なデータが含まれているバッファキャッシュ。データベースインスタンスはメインメモリまたはディスクから読み取る必要があり、バッファキャッシュから読み取るよりも時間がかかるため、パフォーマンスに影響します。

コールドデータ

めったにアクセスされず、通常は過去のデータです。この種類のデータをクエリする場合、通常は低速なクエリでも問題ありません。このデータを低パフォーマンスで安価なストレージ階層またはクラスに移動すると、コストを削減することができます。

コンピュータビジョン (CV)

機械学習を使用してデジタルイメージやビデオなどのビジュアル形式から情報を分析および抽出する [AI](#) の分野。例えば、Amazon SageMaker AI は CV 用の画像処理アルゴリズムを提供します。

設定ドリフト

ワークロードの場合、設定は想定状態から変化します。これにより、ワークロードが非準拠になる可能性があり、通常は段階的かつ意図的ではありません。

構成管理データベース (CMDB)

データベースとその IT 環境 (ハードウェアとソフトウェアの両方のコンポーネントとその設定を含む) に関する情報を保存、管理するリポジトリ。通常、CMDB のデータは、移行のポートフォリオの検出と分析の段階で使用します。

コンフォーマンスパック

コンプライアンスチェックとセキュリティチェックをカスタマイズするためにアセンブルできる AWS Config ルールと修復アクションのコレクション。YAML テンプレートを使用して、コンフォーマンスパックを AWS アカウント および リージョンの単一のエンティティとしてデプロイすることも、組織全体にデプロイすることもできます。詳細については、AWS Config ドキュメントの「[コンフォーマンスパック](#)」を参照してください。

継続的インテグレーションと継続的デリバリー (CI/CD)

ソフトウェアリリースプロセスのソース、ビルド、テスト、ステージング、本番の各ステージを自動化するプロセス。CI/CD は一般的にパイプラインと呼ばれます。プロセスの自動化、生産性の向上、コード品質の向上、配信の加速化を可能にします。詳細については、「[継続的デリバ](#)

[リーの利点](#)」を参照してください。CD は継続的デプロイ (Continuous Deployment) の略語でもあります。詳細については「[継続的デリバリーと継続的なデプロイ](#)」を参照してください。

CV

[「コンピュータビジョン」](#)を参照してください。

D

保管中のデータ

ストレージ内にあるデータなど、常に自社のネットワーク内にあるデータ。

データ分類

ネットワーク内のデータを重要度と機密性に基づいて識別、分類するプロセス。データに適した保護および保持のコントロールを判断する際に役立つため、あらゆるサイバーセキュリティのリスク管理戦略において重要な要素です。データ分類は、AWS Well-Architected フレームワークのセキュリティの柱のコンポーネントです。詳細については、[データ分類](#)を参照してください。

データドリフト

実稼働データと ML モデルのトレーニングに使用されたデータとの間に有意な差異が生じたり、入力データが時間の経過と共に有意に変化したりすることです。データドリフトは、ML モデル予測の全体的な品質、精度、公平性を低下させる可能性があります。

転送中のデータ

ネットワーク内 (ネットワークリソース間など) を活発に移動するデータ。

データメッシュ

一元管理とガバナンスを備えた分散型の分散データ所有権を提供するアーキテクチャフレームワーク。

データ最小化

厳密に必要なデータのみを収集し、処理するという原則。でデータ最小化を実践 AWS クラウドすることで、プライバシーリスク、コスト、分析のカーボンフットプリントを削減できます。

データ境界

AWS 環境内の一連の予防ガードレール。信頼された ID のみが、期待されるネットワークから信頼されたリソースにアクセスできるようにします。詳細については、「[でのデータ境界の構築 AWS](#)」を参照してください。

データの事前処理

raw データをお客様の機械学習モデルで簡単に解析できる形式に変換すること。データの事前処理とは、特定の列または行を削除して、欠落している、矛盾している、または重複する値に対処することを意味します。

データの出所

データの生成、送信、保存の方法など、データのライフサイクル全体を通じてデータの出所と履歴を追跡するプロセス。

データの種類

データを収集、処理している個人。

データウェアハウス

分析などのビジネスインテリジェンスをサポートするデータ管理システム。データウェアハウスには、通常、大量の履歴データが含まれており、クエリや分析に使用されます。

データベース定義言語 (DDL)

データベース内のテーブルやオブジェクトの構造を作成または変更するためのステートメントまたはコマンド。

データベース操作言語 (DML)

データベース内の情報を変更 (挿入、更新、削除) するためのステートメントまたはコマンド。

DDL

[「データベース定義言語」](#)を参照してください。

ディープアンサンブル

予測のために複数の深層学習モデルを組み合わせる。ディープアンサンブルを使用して、より正確な予測を取得したり、予測の不確実性を推定したりできます。

ディープラーニング

人工ニューラルネットワークの複数層を使用して、入力データと対象のターゲット変数の間のマッピングを識別する機械学習サブフィールド。

多層防御

一連のセキュリティメカニズムとコントロールをコンピュータネットワーク全体に層状に重ねて、ネットワークとその内部にあるデータの機密性、整合性、可用性を保護する情報セキュリティ

ティの手法。この戦略を採用するときは AWS、AWS Organizations 構造の異なるレイヤーに複数のコントロールを追加して、リソースの安全性を確保します。たとえば、多層防御アプローチでは、多要素認証、ネットワークセグメンテーション、暗号化を組み合わせることができます。

委任管理者

では AWS Organizations、互換性のあるサービスが AWS メンバーアカウントを登録して組織のアカウントを管理し、そのサービスのアクセス許可を管理できます。このアカウントを、そのサービスの委任管理者と呼びます。詳細、および互換性のあるサービスの一覧は、AWS Organizations ドキュメントの[AWS Organizationsで利用できるサービス](#)を参照してください。

デプロイ

アプリケーション、新機能、コードの修正をターゲットの環境で利用できるようにするプロセス。デプロイでは、コードベースに変更を施した後、アプリケーションの環境でそのコードベースを構築して実行します。

開発環境

[???「環境」](#)を参照してください。

検出管理

イベントが発生したときに、検出、ログ記録、警告を行うように設計されたセキュリティコントロール。これらのコントロールは副次的な防衛手段であり、実行中の予防的コントロールをすり抜けたセキュリティイベントをユーザーに警告します。詳細については、Implementing security controls on AWSの[Detective controls](#)を参照してください。

開発バリューストリームマッピング (DVSM)

ソフトウェア開発ライフサイクルのスピードと品質に悪影響を及ぼす制約を特定し、優先順位を付けるために使用されるプロセス。DVSM は、もともとリーンマニファクチャリング・プラクティスのために設計されたバリューストリームマッピング・プロセスを拡張したものです。ソフトウェア開発プロセスを通じて価値を創造し、動かすために必要なステップとチームに焦点を当てています。

デジタルツイン

建物、工場、産業機器、生産ラインなど、現実世界のシステムを仮想的に表現したものです。デジタルツインは、予知保全、リモートモニタリング、生産最適化をサポートします。

ディメンションテーブル

[スタースキーマ](#)では、ファクトテーブル内の量的データに関するデータ属性を含む小さなテーブル。ディメンションテーブル属性は通常、テキストフィールドまたはテキストのように動作する

離散数値です。これらの属性は、クエリの制約、フィルタリング、結果セットのラベル付けに一般的に使用されます。

ディザスタ

ワークロードまたはシステムが、導入されている主要な場所でのビジネス目標の達成を妨げるイベント。これらのイベントは、自然災害、技術的障害、または意図しない設定ミスやマルウェア攻撃などの人間の行動の結果である場合があります。

ディザスタリカバリ (DR)

[災害](#)によるダウンタイムとデータ損失を最小限に抑えるために使用する戦略とプロセス。詳細については、AWS Well-Architected フレームワークの「[でのワークロードのディザスタリカバリ](#) [AWS: クラウドでのリカバリ](#)」を参照してください。

DML

「[データベース操作言語](#)」を参照してください。

ドメイン駆動型設計

各コンポーネントが提供している変化を続けるドメイン、またはコアビジネス目標にコンポーネントを接続して、複雑なソフトウェアシステムを開発するアプローチ。この概念は、エリック・エヴァンスの著書、Domain-Driven Design: Tackling Complexity in the Heart of Software (ドメイン駆動設計:ソフトウェアの中心における複雑さへの取り組み) で紹介されています (ボストン: Addison-Wesley Professional、2003)。strangler fig パターンでドメイン駆動型設計を使用する方法の詳細については、[コンテナと Amazon API Gateway を使用して、従来の Microsoft ASP.NET \(ASMX\) ウェブサービスを段階的にモダナイズ](#) を参照してください。

DR

「[ディザスタリカバリ](#)」を参照してください。

ドリフト検出

ベースライン設定からの偏差を追跡します。例えば、AWS CloudFormation を使用して[システムリソースのドリフトを検出](#)したり、を使用して AWS Control Tower、ガバナンス要件への準拠に影響する[ランディングゾーンの変更を検出](#)したりできます。

DVSM

「[開発値ストリームマッピング](#)」を参照してください。

E

EDA

[「探索的データ分析」](#)を参照してください。

EDI

[「電子データ交換」](#)を参照してください。

エッジコンピューティング

IoT ネットワークのエッジにあるスマートデバイスの計算能力を高めるテクノロジー。[クラウドコンピューティング](#)と比較すると、エッジコンピューティングは通信レイテンシーを短縮し、応答時間を短縮できます。

電子データ交換 (EDI)

組織間のビジネスドキュメントの自動交換。詳細については、[「電子データ交換とは」](#)を参照してください。

暗号化

人間が読み取り可能なプレーンテキストデータを暗号文に変換するコンピューティングプロセス。

暗号化キー

暗号化アルゴリズムが生成した、ランダム化されたビットからなる暗号文字列。キーの長さは決まっておらず、各キーは予測できないように、一意になるように設計されています。

エンディアン

コンピュータメモリにバイトが格納される順序。ビッグエンディアンシステムでは、最上位バイトが最初に格納されます。リトルエンディアンシステムでは、最下位バイトが最初に格納されます。

エンドポイント

[「サービスエンドポイント」](#)を参照してください。

エンドポイントサービス

仮想プライベートクラウド (VPC) 内でホストして、他のユーザーと共有できるサービス。を使用してエンドポイントサービスを作成し AWS PrivateLink、他の AWS アカウント または AWS Identity and Access Management (IAM) プリンシパルにアクセス許可を付与できます。これら

のアカウントまたはプリンシパルは、インターフェイス VPC エンドポイントを作成することで、エンドポイントサービスにプライベートに接続できます。詳細については、Amazon Virtual Private Cloud (Amazon VPC) ドキュメントの「[エンドポイントサービスを作成する](#)」を参照してください。

エンタープライズリソースプランニング (ERP)

エンタープライズの主要なビジネスプロセス (会計、[MES](#)、プロジェクト管理など) を自動化および管理するシステム。

エンベロープ暗号化

暗号化キーを、別の暗号化キーを使用して暗号化するプロセス。詳細については、AWS Key Management Service (AWS KMS) ドキュメントの「[エンベロープ暗号化](#)」を参照してください。

環境

実行中のアプリケーションのインスタンス。クラウドコンピューティングにおける一般的な環境の種類は以下のとおりです。

- 開発環境 — アプリケーションのメンテナンスを担当するコアチームのみが利用できる、実行中のアプリケーションのインスタンス。開発環境は、上位の環境に昇格させる変更をテストするときに使用します。このタイプの環境は、テスト環境と呼ばれることもあります。
- 下位環境 — 初期ビルドやテストに使用される環境など、アプリケーションのすべての開発環境。
- 本番環境 — エンドユーザーがアクセスできる、実行中のアプリケーションのインスタンス。CI/CD パイプラインでは、本番環境が最後のデプロイ環境になります。
- 上位環境 — コア開発チーム以外のユーザーがアクセスできるすべての環境。これには、本番環境、本番前環境、ユーザー承認テスト環境などが含まれます。

エピック

アジャイル方法論で、お客様の作業の整理と優先順位付けに役立つ機能力カテゴリー。エピックでは、要件と実装タスクの概要についてハイレベルな説明を提供します。例えば、AWS CAF セキュリティエピックには、ID とアクセスの管理、検出コントロール、インフラストラクチャセキュリティ、データ保護、インシデント対応が含まれます。AWS 移行戦略のエピックの詳細については、[プログラム実装ガイド](#) を参照してください。

ERP

「[エンタープライズリソース計画](#)」を参照してください。

探索的データ分析 (EDA)

データセットを分析してその主な特性を理解するプロセス。お客様は、データを収集または集計してから、パターンの検出、異常の検出、および前提条件のチェックのための初期調査を実行します。EDA は、統計の概要を計算し、データの可視化を作成することによって実行されます。

F

ファクトテーブル

[星スキーマ](#)の中央テーブル。事業運営に関する量的データを保存します。通常、ファクトテーブルには、メジャーを含む列とディメンションテーブルへの外部キーを含む列の 2 つのタイプの列が含まれます。

フェイルファスト

開発ライフサイクルを短縮するために頻繁で段階的なテストを使用する哲学。これはアジャイルアプローチの重要な部分です。

障害分離の境界

では AWS クラウド、アベイラビリティーゾーン AWS リージョン、コントロールプレーン、データプレーンなどの境界で、障害の影響を制限し、ワークロードの耐障害性を向上させるのに役立ちます。詳細については、[AWS 「障害分離境界」](#)を参照してください。

機能ブランチ

[「ブランチ」](#)を参照してください。

特徴量

お客様が予測に使用する入力データ。例えば、製造コンテキストでは、特徴量は製造ラインから定期的にキャプチャされるイメージの可能性もあります。

特徴量重要度

モデルの予測に対する特徴量の重要性。これは通常、Shapley Additive Deskonations (SHAP) や積分勾配など、さまざまな手法で計算できる数値スコアで表されます。詳細については、[「を使用した機械学習モデルの解釈可能性 AWS」](#)を参照してください。

機能変換

追加のソースによるデータのエンリッチ化、値のスケーリング、単一のデータフィールドからの複数の情報セットの抽出など、機械学習プロセスのデータを最適化すること。これにより、機械

学習モデルはデータの恩恵を受けることができます。例えば、「2021-05-27 00:15:37」の日付を「2021 年」、「5 月」、「木」、「15」に分解すると、学習アルゴリズムがさまざまなデータコンポーネントに関連する微妙に異なるパターンを学習するのに役立ちます。

数ショットプロンプト

同様のタスクの実行を求める前に、タスクと必要な出力を示す少数の例を [LLM](#) に提供します。この手法は、プロンプトに埋め込まれた例 (ショット) からモデルが学習するコンテキスト内学習のアプリケーションです。少数ショットプロンプトは、特定のフォーマット、推論、またはドメインの知識を必要とするタスクに効果的です。[「ゼロショットプロンプト」](#) も参照してください。

FGAC

[「きめ細かなアクセスコントロール」](#) を参照してください。

きめ細かなアクセス制御 (FGAC)

複数の条件を使用してアクセス要求を許可または拒否すること。

フラッシュカット移行

段階的なアプローチを使用する代わりに、[変更データキャプチャ](#) による継続的なデータレプリケーションを使用して、可能な限り短時間でデータを移行するデータベース移行方法。目的はダウンタイムを最小限に抑えることです。

FM

[「基盤モデル」](#) を参照してください。

基盤モデル (FM)

一般化データとラベル付けされていないデータの大規模なデータセットでトレーニングされている大規模な深層学習ニューラルネットワーク。FMs は、言語の理解、テキストと画像の生成、自然言語の会話など、さまざまな一般的なタスクを実行できます。詳細については、[「基盤モデルとは」](#) を参照してください。

G

生成 AI

大量のデータでトレーニングされ、シンプルなテキストプロンプトを使用してイメージ、動画、テキスト、オーディオなどの新しいコンテンツやアーティファクトを作成できる [AI](#) モデルのサブセット。詳細については、[「生成 AI とは」](#) を参照してください。

ジオブロッキング

[地理的制限](#)を参照してください。

地理的制限 (ジオブロッキング)

特定の国のユーザーがコンテンツ配信にアクセスできないようにするための、Amazon CloudFront のオプション。アクセスを許可する国と禁止する国は、許可リストまたは禁止リストを使って指定します。詳細については、CloudFront ドキュメントの[コンテンツの地理的ディストリビューションの制限](#)を参照してください。

Gitflow ワークフロー

下位環境と上位環境が、ソースコードリポジトリでそれぞれ異なるブランチを使用する方法。Gitflow ワークフローはレガシーと見なされ、[トランクベースのワークフロー](#)はモダンで推奨されるアプローチです。

ゴールデンイメージ

そのシステムまたはソフトウェアの新しいインスタンスをデプロイするためのテンプレートとして使用されるシステムまたはソフトウェアのスナップショット。例えば、製造では、ゴールデンイメージを使用して複数のデバイスにソフトウェアをプロビジョニングし、デバイス製造オペレーションの速度、スケーラビリティ、生産性を向上させることができます。

グリーンフィールド戦略

新しい環境に既存のインフラストラクチャが存在しないこと。システムアーキテクチャにグリーンフィールド戦略を導入する場合、既存のインフラストラクチャ (別名[ブラウンフィールド](#)) との互換性の制約を受けることなく、あらゆる新しいテクノロジーを選択できます。既存のインフラストラクチャを拡張している場合は、ブラウンフィールド戦略とグリーンフィールド戦略を融合させることもできます。

ガードレール

組織単位 (OU) 全般のリソース、ポリシー、コンプライアンスを管理するのに役立つ概略的なルール。予防ガードレールは、コンプライアンス基準に一致するようにポリシーを実施します。これらは、サービスコントロールポリシーと IAM アクセス許可の境界を使用して実装されます。検出ガードレールは、ポリシー違反やコンプライアンス上の問題を検出し、修復のためのアラートを発信します。これらは AWS Config、Amazon GuardDuty AWS Security Hub、AWS Trusted Advisor Amazon Inspector、およびカスタム AWS Lambda チェックを使用して実装されます。

H

HA

[「高可用性」](#)を参照してください。

異種混在データベースの移行

別のデータベースエンジンを使用するターゲットデータベースへお客様の出典データベースの移行 (例えば、Oracle から Amazon Aurora)。異種間移行は通常、アーキテクチャの再設計作業の一部であり、スキーマの変換は複雑なタスクになる可能性があります。[AWS は、スキーマの変換に役立つ AWS SCTを提供します。](#)

ハイアベイラビリティ (HA)

課題や災害が発生した場合に、介入なしにワークロードを継続的に運用できること。HA システムは、自動的にフェイルオーバーし、一貫して高品質のパフォーマンスを提供し、パフォーマンスへの影響を最小限に抑えながらさまざまな負荷や障害を処理するように設計されています。

ヒストリアンのモダナイゼーション

製造業のニーズによりよく応えるために、オペレーションテクノロジー (OT) システムをモダナイズし、アップグレードするためのアプローチ。ヒストリアンは、工場内のさまざまなソースからデータを収集して保存するために使用されるデータベースの一種です。

ホールドアウトデータ

[機械学習](#)モデルのトレーニングに使用されるデータセットから保留される、ラベル付きの履歴データの一部。モデル予測をホールドアウトデータと比較することで、ホールドアウトデータを使用してモデルのパフォーマンスを評価できます。

同種データベースの移行

お客様の出典データベースを、同じデータベースエンジンを共有するターゲットデータベース (Microsoft SQL Server から Amazon RDS for SQL Server など) に移行する。同種間移行は、通常、リホストまたはリプラットフォーム化の作業の一部です。ネイティブデータベースユーティリティを使用して、スキーマを移行できます。

ホットデータ

リアルタイムデータや最近の翻訳データなど、頻繁にアクセスされるデータ。通常、このデータには高速なクエリ応答を提供する高性能なストレージ階層またはクラスが必要です。

ホットフィックス

本番環境の重大な問題を修正するために緊急で配布されるプログラム。緊急性が高いため、通常の DevOps のリリースワークフローからは外れた形で実施されます。

ハイパーケア期間

カットオーバー直後、移行したアプリケーションを移行チームがクラウドで管理、監視して問題に対処する期間。通常、この期間は 1~4 日です。ハイパーケア期間が終了すると、アプリケーションに対する責任は一般的に移行チームからクラウドオペレーションチームに移ります。

I

IaC

[「Infrastructure as Code」](#) を参照してください。

ID ベースのポリシー

AWS クラウド 環境内のアクセス許可を定義する 1 つ以上の IAM プリンシパルにアタッチされたポリシー。

アイドル状態のアプリケーション

90 日間の平均的な CPU およびメモリ使用率が 5~20% のアプリケーション。移行プロジェクトでは、これらのアプリケーションを廃止するか、オンプレミスに保持するのが一般的です。

IIoT

[「産業用モノのインターネット」](#) を参照してください。

イミュータブルインフラストラクチャ

既存のインフラストラクチャを更新、パッチ適用、または変更する代わりに、本番環境のワークロード用に新しいインフラストラクチャをデプロイするモデル。イミュータブルインフラストラクチャは、本質的に [ミュータブルインフラストラクチャ](#) よりも一貫性、信頼性、予測性が高くなります。詳細については、AWS 「Well-Architected フレームワーク」の [「イミュータブルインフラストラクチャを使用したデプロイ」](#) のベストプラクティスを参照してください。

インバウンド (受信) VPC

AWS マルチアカウントアーキテクチャでは、アプリケーションの外部からネットワーク接続を受け入れ、検査し、ルーティングする VPC。 [AWS Security Reference Architecture](#) では、アプリ

ケーションとより広範なインターネット間の双方向のインターフェイスを保護するために、インバウンド、アウトバウンド、インスペクションの各 VPC を使用してネットワークアカウントを設定することを推奨しています。

増分移行

アプリケーションを 1 回ですべてカットオーバーするのではなく、小さい要素に分けて移行するカットオーバー戦略。例えば、最初は少数のマイクロサービスまたはユーザーのみを新しいシステムに移行する場合があります。すべてが正常に機能することを確認できたら、残りのマイクロサービスやユーザーを段階的に移行し、レガシーシステムを廃止できるようにします。この戦略により、大規模な移行に伴うリスクが軽減されます。

インダストリー 4.0

2016 年に [Klaus Schwab](#) によって導入された用語で、接続、リアルタイムデータ、オートメーション、分析、AI/ML の進歩によるビジネスプロセスのモダナイゼーションを指します。

インフラストラクチャ

アプリケーションの環境に含まれるすべてのリソースとアセット。

Infrastructure as Code (IaC)

アプリケーションのインフラストラクチャを一連の設定ファイルを使用してプロビジョニングし、管理するプロセス。IaC は、新しい環境を再現可能で信頼性が高く、一貫性のあるものにするため、インフラストラクチャを一元的に管理し、リソースを標準化し、スケールを迅速に行えるように設計されています。

産業分野における IoT (IIoT)

製造、エネルギー、自動車、ヘルスケア、ライフサイエンス、農業などの産業部門におけるインターネットに接続されたセンサーやデバイスの使用。詳細については、「[Building an industrial Internet of Things \(IIoT\) digital transformation strategy](#)」を参照してください。

インスペクション VPC

AWS マルチアカウントアーキテクチャでは、VPC (同一または異なる AWS リージョン)、インターネット、オンプレミスネットワーク間のネットワークトラフィックの検査を管理する一元化された VPCs。 [AWS Security Reference Architecture](#) では、アプリケーションとより広範なインターネット間の双方向のインターフェイスを保護するために、インバウンド、アウトバウンド、インスペクションの各 VPC を使用してネットワークアカウントを設定することを推奨しています。

IoT

インターネットまたはローカル通信ネットワークを介して他のデバイスやシステムと通信する、センサーまたはプロセッサが組み込まれた接続済み物理オブジェクトのネットワーク。詳細については、「[IoT とは](#)」を参照してください。

解釈可能性

機械学習モデルの特性で、モデルの予測がその入力にどのように依存するかを人間が理解できる度合いを表します。詳細については、「[を使用した機械学習モデルの解釈可能性 AWS](#)」を参照してください。

IoT

「[モノのインターネット](#)」を参照してください。

IT 情報ライブラリ (ITIL)

IT サービスを提供し、これらのサービスをビジネス要件に合わせるための一連のベストプラクティス。ITIL は ITSM の基盤を提供します。

IT サービス管理 (ITSM)

組織の IT サービスの設計、実装、管理、およびサポートに関連する活動。クラウドオペレーションと ITSM ツールの統合については、[オペレーション統合ガイド](#) を参照してください。

ITIL

「[IT 情報ライブラリ](#)」を参照してください。

ITSM

「[IT サービス管理](#)」を参照してください。

L

ラベルベースアクセス制御 (LBAC)

強制アクセス制御 (MAC) の実装で、ユーザーとデータ自体にそれぞれセキュリティラベル値が明示的に割り当てられます。ユーザーセキュリティラベルとデータセキュリティラベルが交差する部分によって、ユーザーに表示される行と列が決まります。

ランディングゾーン

ランディングゾーンは、スケーラブルで安全な、適切に設計されたマルチアカウント AWS 環境です。これは、組織がセキュリティおよびインフラストラクチャ環境に自信を持ってワークロー

ドとアプリケーションを迅速に起動してデプロイできる出発点です。ランディングゾーンの詳細については、[安全でスケーラブルなマルチアカウント AWS 環境のセットアップ](#) を参照してください。

大規模言語モデル (LLM)

大量のデータに対して事前トレーニングされた深層学習 [AI](#) モデル。LLM は、質問への回答、ドキュメントの要約、テキストの他の言語への翻訳、文の完了など、複数のタスクを実行できます。詳細については、[LLMs](#) を参照してください。

大規模な移行

300 台以上のサーバの移行。

LBAC

[「ラベルベースのアクセスコントロール」](#) を参照してください。

最小特権

タスクの実行には必要最低限の権限を付与するという、セキュリティのベストプラクティス。詳細については、IAM ドキュメントの [最小特権アクセス許可を適用する](#) を参照してください。

リフトアンドシフト

[「7 Rs」](#) を参照してください。

リトルエンディアンシステム

最下位バイトを最初に格納するシステム。 [エンディアン性](#) も参照してください。

LLM

[「大規模言語モデル」](#) を参照してください。

下位環境

[「環境」](#) を参照してください。

M

機械学習 (ML)

パターン認識と学習にアルゴリズムと手法を使用する人工知能の一種。ML は、モノのインターネット (IoT) データなどの記録されたデータを分析して学習し、パターンに基づく統計モデルを生成します。詳細については、[「機械学習」](#) を参照してください。

メインランチ

[「ブランチ」](#)を参照してください。

マルウェア

コンピュータのセキュリティまたはプライバシーを侵害するように設計されたソフトウェア。マルウェアは、コンピュータシステムの中断、機密情報の漏洩、不正アクセスにつながる可能性があります。マルウェアの例としては、ウイルス、ワーム、ランサムウェア、トロイの木馬、スパイウェア、キーロガーなどがあります。

マネージドサービス

AWS のサービス はインフラストラクチャレイヤー、オペレーティングシステム、プラットフォーム AWS を運用し、エンドポイントにアクセスしてデータを保存および取得します。Amazon Simple Storage Service (Amazon S3) と Amazon DynamoDB は、マネージドサービスの例です。これらは抽象化されたサービスとも呼ばれます。

製造実行システム (MES)

生産プロセスを追跡、モニタリング、文書化、制御するためのソフトウェアシステムで、原材料を工場の完成製品に変換します。

MAP

[「移行促進プログラム」](#)を参照してください。

メカニズム

ツールを作成し、ツールの導入を推進し、調整を行うために結果を検査する完全なプロセス。メカニズムは、動作中にそれ自体を強化して改善するサイクルです。詳細については、AWS「Well-Architected フレームワーク」の[「メカニズムの構築」](#)を参照してください。

メンバーアカウント

組織の一部である管理アカウント AWS アカウント 以外のすべて AWS Organizations。アカウントが組織のメンバーになることができるのは、一度に 1 つのみです。

MES

[「製造実行システム」](#)を参照してください。

メッセージキューイングテレメトリトランスポート (MQTT)

リソースに制約のある [IoT](#) デバイス用の、[パブリッシュ/サブスクライブ](#)パターンに基づく軽量 machine-to-machine (M2M) 通信プロトコル。

マイクロサービス

明確に定義された API を介して通信し、通常は小規模な自己完結型のチームが所有する、小規模で独立したサービスです。例えば、保険システムには、販売やマーケティングなどのビジネス機能、または購買、請求、分析などのサブドメインにマッピングするマイクロサービスが含まれる場合があります。マイクロサービスの利点には、俊敏性、柔軟なスケーリング、容易なデプロイ、再利用可能なコード、回復力などがあります。詳細については、[AWS「サーバーレスサービスを使用したマイクロサービスの統合」](#)を参照してください。

マイクロサービスアーキテクチャ

各アプリケーションプロセスをマイクロサービスとして実行する独立したコンポーネントを使用してアプリケーションを構築するアプローチ。これらのマイクロサービスは、軽量 API を使用して、明確に定義されたインターフェイスを介して通信します。このアーキテクチャの各マイクロサービスは、アプリケーションの特定の機能に対する需要を満たすように更新、デプロイ、およびスケーリングできます。詳細については、「[でのマイクロサービスの実装 AWS](#)」を参照してください。

Migration Acceleration Program (MAP)

組織がクラウドに移行するための強力な運用基盤を構築し、移行の初期コストを相殺するのに役立つコンサルティングサポート、トレーニング、サービスを提供する AWS プログラム。MAP には、組織的な方法でレガシー移行を実行するための移行方法論と、一般的な移行シナリオを自動化および高速化する一連のツールが含まれています。

大規模な移行

アプリケーションポートフォリオの大部分を次々にクラウドに移行し、各ウェーブでより多くのアプリケーションを高速に移動させるプロセス。この段階では、以前の段階から学んだベストプラクティスと教訓を使用して、移行ファクトリー チーム、ツール、プロセスのうち、オートメーションとアジャイルデリバリーによってワークロードの移行を合理化します。これは、[AWS 移行戦略](#) の第 3 段階です。

移行ファクトリー

自動化された俊敏性のあるアプローチにより、ワークロードの移行を合理化する部門横断的なチーム。移行ファクトリーチームには、通常、運用、ビジネスアナリストおよび所有者、移行エンジニア、デベロッパー、およびスプリントで作業する DevOps プロフェッショナルが含まれます。エンタープライズアプリケーションポートフォリオの 20～50% は、ファクトリーのアプローチによって最適化できる反復パターンで構成されています。詳細については、このコンテンツセットの[移行ファクトリーに関する解説](#)と [Cloud Migration Factory ガイド](#)を参照してください。

移行メタデータ

移行を完了するために必要なアプリケーションおよびサーバーに関する情報。移行パターンごとに、異なる一連の移行メタデータが必要です。移行メタデータの例としては、ターゲットサブネット、セキュリティグループ、AWS アカウントなどがあります。

移行パターン

移行戦略、移行先、および使用する移行アプリケーションまたはサービスを詳述する、反復可能な移行タスク。例: AWS Application Migration Service を使用して Amazon EC2 への移行をリホストします。

Migration Portfolio Assessment (MPA)

に移行するためのビジネスケースを検証するための情報を提供するオンラインツール AWS クラウド。MPA は、詳細なポートフォリオ評価 (サーバーの適切なサイジング、価格設定、TCO 比較、移行コスト分析) および移行プラン (アプリケーションデータの分析とデータ収集、アプリケーションのグループ化、移行の優先順位付け、およびウェーブプランニング) を提供します。[MPA ツール](#) (ログインが必要) は、すべての AWS コンサルタントと APN パートナーコンサルタントが無料で利用できます。

移行準備状況評価 (MRA)

AWS CAF を使用して、組織のクラウド準備状況に関するインサイトを取得し、長所と短所を特定し、特定されたギャップを埋めるためのアクションプランを構築するプロセス。詳細については、[移行準備状況ガイド](#) を参照してください。MRA は、[AWS 移行戦略](#)の第一段階です。

移行戦略

ワークロードを に移行するために使用するアプローチ AWS クラウド。詳細については、この用語集の「[7 Rs](#) エントリ」と「[組織を動員して大規模な移行を加速する](#)」を参照してください。

ML

[??? 「機械学習」](#) を参照してください。

モダナイゼーション

古い (レガシーまたはモノリシック) アプリケーションとそのインフラストラクチャをクラウド内の俊敏で弾力性のある高可用性システムに変換して、コストを削減し、効率を高め、イノベーションを活用します。詳細については、「」の「[アプリケーションをモダナイズするための戦略 AWS クラウド](#)」を参照してください。

モダナイゼーション準備状況評価

組織のアプリケーションのモダナイゼーションの準備状況を判断し、利点、リスク、依存関係を特定し、組織がこれらのアプリケーションの将来の状態をどの程度適切にサポートできるかを決定するのに役立つ評価。評価の結果として、ターゲットアーキテクチャのブループリント、モダナイゼーションプロセスの開発段階とマイルストーンを詳述したロードマップ、特定されたギャップに対処するためのアクションプランが得られます。詳細については、[『』の「アプリケーションのモダナイゼーション準備状況の評価 AWS クラウド」](#)を参照してください。

モノリシックアプリケーション (モノリス)

緊密に結合されたプロセスを持つ単一のサービスとして実行されるアプリケーション。モノリシックアプリケーションにはいくつかの欠点があります。1つのアプリケーション機能エクスペリエンスの需要が急増する場合は、アーキテクチャ全体をスケーリングする必要があります。モノリシックアプリケーションの特徴を追加または改善することは、コードベースが大きくなると複雑になります。これらの問題に対処するには、マイクロサービスアーキテクチャを使用できます。詳細については、[モノリスをマイクロサービスに分解する](#)を参照してください。

MPA

[「移行ポートフォリオ評価」](#)を参照してください。

MQTT

[「Message Queuing Telemetry Transport」](#)を参照してください。

多クラス分類

複数のクラスの予測を生成するプロセス (2 つ以上の結果の 1 つを予測します)。例えば、機械学習モデルが、「この製品は書籍、自動車、電話のいずれですか?」または、「このお客様にとって最も関心のある商品のカテゴリはどれですか?」と聞くかもしれません。

ミュータブルインフラストラクチャ

本番ワークロードの既存のインフラストラクチャを更新および変更するモデル。Well-Architected AWS フレームワークでは、一貫性、信頼性、予測可能性を向上させるために、[イミュータブルインフラストラクチャ](#)の使用をベストプラクティスとして推奨しています。

O

OAC

[「オリジンアクセスコントロール」](#)を参照してください。

OAI

[「オリジンアクセスアイデンティティ」](#)を参照してください。

OCM

[「組織変更管理」](#)を参照してください。

オフライン移行

移行プロセス中にソースワークロードを停止させる移行方法。この方法はダウンタイムが長くなるため、通常は重要ではない小規模なワークロードに使用されます。

OI

[「オペレーションの統合」](#)を参照してください。

OLA

[「運用レベルの契約」](#)を参照してください。

オンライン移行

ソースワークロードをオフラインにせずにターゲットシステムにコピーする移行方法。ワークロードに接続されているアプリケーションは、移行中も動作し続けることができます。この方法はダウンタイムがゼロから最小限で済むため、通常は重要な本番稼働環境のワークロードに使用されます。

OPC-UA

[「Open Process Communications - Unified Architecture」](#)を参照してください。

オープンプロセス通信 - 統合アーキテクチャ (OPC-UA)

産業用オートメーション用のmachine-to-machine (M2M) 通信プロトコル。OPC-UA は、データの暗号化、認証、認可スキームを備えた相互運用性標準を提供します。

オペレーショナルレベルアグリーメント (OLA)

サービスレベルアグリーメント (SLA) をサポートするために、どの機能的 IT グループが互いに提供することを約束するかを明確にする契約。

運用準備状況レビュー (ORR)

インシデントや潜在的な障害の理解、評価、防止、または範囲の縮小に役立つ質問とそれに関連するベストプラクティスのチェックリスト。詳細については、AWS Well-Architected フレームワークの [「Operational Readiness Reviews \(ORR\)」](#)を参照してください。

運用テクノロジー (OT)

産業オペレーション、機器、インフラストラクチャを制御するために物理環境と連携するハードウェアおよびソフトウェアシステム。製造では、OT と情報技術 (IT) システムの統合が、[Industry 4.0](#) 変換の主な焦点です。

オペレーション統合 (OI)

クラウドでオペレーションをモダナイズするプロセスには、準備計画、オートメーション、統合が含まれます。詳細については、[オペレーション統合ガイド](#) を参照してください。

組織の証跡

組織 AWS アカウント 内のすべてのイベント AWS CloudTrail をログに記録する、によって作成された証跡 AWS Organizations。証跡は、組織に含まれている各 AWS アカウントに作成され、各アカウントのアクティビティを追跡します。詳細については、CloudTrail ドキュメントの[組織の証跡の作成](#)を参照してください。

組織変更管理 (OCM)

人材、文化、リーダーシップの観点から、主要な破壊的なビジネス変革を管理するためのフレームワーク。OCM は、変化の導入を加速し、移行問題に対処し、文化や組織の変化を推進することで、組織が新しいシステムと戦略の準備と移行するのを支援します。AWS 移行戦略では、クラウド導入プロジェクトに必要な変化のスピードのため、このフレームワークは人材アクセラレーションと呼ばれます。詳細については、[OCM ガイド](#) を参照してください。

オリジンアクセスコントロール (OAC)

Amazon Simple Storage Service (Amazon S3) コンテンツを保護するための、CloudFront のアクセス制限の強化オプション。OAC は AWS リージョン、すべての S3 バケット、AWS KMS (SSE-KMS) によるサーバー側の暗号化、S3 バケットへの動的 PUT および DELETE リクエストをサポートします。

オリジンアクセスアイデンティティ (OAI)

CloudFront の、Amazon S3 コンテンツを保護するためのアクセス制限オプション。OAI を使用すると、CloudFront が、Amazon S3 に認証可能なプリンシパルを作成します。認証されたプリンシパルは、S3 バケット内のコンテンツに、特定の CloudFront ディストリビューションを介してのみアクセスできます。[OAC](#) も併せて参照してください。OAC では、より詳細な、強化されたアクセスコントロールが可能です。

ORR

[「運用準備状況レビュー」](#) を参照してください。

OT

[「運用テクノロジー」](#)を参照してください。

アウトバウンド (送信) VPC

AWS マルチアカウントアーキテクチャでは、アプリケーション内から開始されたネットワーク接続を処理する VPC。[AWS Security Reference Architecture](#) では、アプリケーションとより広範なインターネット間の双方向のインターフェイスを保護するために、インバウンド、アウトバウンド、インスペクションの各 VPC を使用してネットワークアカウントを設定することを推奨しています。

P

アクセス許可の境界

ユーザーまたはロールが使用できるアクセス許可の上限を設定する、IAM プリンシパルにアタッチされる IAM 管理ポリシー。詳細については、IAM ドキュメントの[アクセス許可の境界](#)を参照してください。

個人を特定できる情報 (PII)

直接閲覧した場合、または他の関連データと組み合わせた場合に、個人の身元を合理的に推測するために使用できる情報。PII の例には、氏名、住所、連絡先情報などがあります。

PII

[個人を特定できる情報](#)を参照してください。

プレイブック

クラウドでのコアオペレーション機能の提供など、移行に関連する作業を取り込む、事前定義された一連のステップ。プレイブックは、スクリプト、自動ランブック、またはお客様のモダナイズされた環境を運用するために必要なプロセスや手順の要約などの形式をとることができます。

PLC

[「プログラム可能なロジックコントローラー」](#)を参照してください。

PLM

[「製品ライフサイクル管理」](#)を参照してください。

ポリシー

アクセス許可を定義 ([アイデンティティベースのポリシー](#)を参照)、アクセス条件を指定 ([リソースベースのポリシー](#)を参照)、または の組織内のすべてのアカウントに対する最大アクセス許可を定義 AWS Organizations ([サービスコントロールポリシー](#)を参照) できるオブジェクト。

多言語の永続性

データアクセスパターンやその他の要件に基づいて、マイクロサービスのデータストレージテクノロジーを個別に選択します。マイクロサービスが同じデータストレージテクノロジーを使用している場合、実装上の問題が発生したり、パフォーマンスが低下する可能性があります。マイクロサービスは、要件に最も適合したデータストアを使用すると、より簡単に実装でき、パフォーマンスとスケーラビリティが向上します。詳細については、[マイクロサービスでのデータ永続性の有効化](#) を参照してください。

ポートフォリオ評価

移行を計画するために、アプリケーションポートフォリオの検出、分析、優先順位付けを行うプロセス。詳細については、「[移行準備状況ガイド](#)」を参照してください。

述語

true または を返すクエリ条件。一般的に false は WHERE 句にあります。

述語プッシュダウン

転送前にクエリ内のデータをフィルタリングするデータベースクエリ最適化手法。これにより、リレーショナルデータベースから取得して処理する必要があるデータの量が減少し、クエリのパフォーマンスが向上します。

予防的コントロール

イベントの発生を防ぐように設計されたセキュリティコントロール。このコントロールは、ネットワークへの不正アクセスや好ましくない変更を防ぐ最前線の防御です。詳細については、Implementing security controls on AWSの[Preventative controls](#)を参照してください。

プリンシパル

アクションを実行し AWS、リソースにアクセスできる のエンティティ。このエンティティは通常、IAM AWS アカウントロール、または ユーザーのルートユーザーです。詳細については、IAM ドキュメントの[ロールに関する用語と概念](#)内にあるプリンシパルを参照してください。

プライバシーバイデザイン

開発プロセス全体を通じてプライバシーを考慮するシステムエンジニアリングアプローチ。

プライベートホストゾーン

1 つ以上の VPC 内のドメインとそのサブドメインへの DNS クエリに対し、Amazon Route 53 がどのように応答するかに関する情報を保持するコンテナ。詳細については、Route 53 ドキュメントの「[プライベートホストゾーンの使用](#)」を参照してください。

プロアクティブコントロール

非準拠リソースのデプロイを防ぐように設計された[セキュリティコントロール](#)。これらのコントロールは、プロビジョニング前にリソースをスキャンします。リソースがコントロールに準拠していない場合、プロビジョニングされません。詳細については、AWS Control Tower ドキュメントの「[コントロールリファレンスガイド](#)」および「[セキュリティコントロールの実装](#)」の「[プロアクティブコントロール](#)」を参照してください。 AWS

製品ライフサイクル管理 (PLM)

設計、開発、発売から成長と成熟まで、製品のデータとプロセスのライフサイクル全体にわたる管理。

本番環境

[「環境」](#)を参照してください。

プログラム可能なロジックコントローラー (PLC)

製造では、マシンをモニタリングし、製造プロセスを自動化する、信頼性の高い適応可能なコンピュータです。

プロンプトの連鎖

1 つの [LLM](#) プロンプトの出力を次のプロンプトの入力として使用して、より良いレスポンスを生成します。この手法は、複雑なタスクをサブタスクに分割したり、事前レスポンスを繰り返し改善または拡張したりするために使用されます。これにより、モデルのレスポンスの精度と関連性が向上し、より詳細でパーソナライズされた結果が得られます。

仮名化

データセット内の個人識別子をプレースホルダー値に置き換えるプロセス。仮名化は個人のプライバシー保護に役立ちます。仮名化されたデータは、依然として個人データとみなされます。

パブリッシュ/サブスクライブ (pub/sub)

マイクロサービス間の非同期通信を可能にするパターン。スケーラビリティと応答性を向上させます。たとえば、マイクロサービスベースの [MES](#) では、マイクロサービスは他のマイクロサー

ビスがサブスクライブできるチャンネルにイベントメッセージを発行できます。システムは、公開サービスを変更せずに新しいマイクロサービスを追加できます。

Q

クエリプラン

SQL リレーショナルデータベースシステムのデータにアクセスするために使用される手順などの一連のステップ。

クエリプランのリグレッション

データベースサービスのオプティマイザーが、データベース環境に特定の変更が加えられる前に選択されたプランよりも最適性の低いプランを選択すること。これは、統計、制限事項、環境設定、クエリパラメータのバインディングの変更、およびデータベースエンジンの更新などが原因である可能性があります。

R

RACI マトリックス

[責任、説明責任、相談、通知 \(RACI\)](#) を参照してください。

RAG

[「取得拡張生成」](#) を参照してください。

ランサムウェア

決済が完了するまでコンピュータシステムまたはデータへのアクセスをブロックするように設計された、悪意のあるソフトウェア。

RASCI マトリックス

[責任、説明責任、相談、情報 \(RACI\)](#) を参照してください。

RCAC

[「行と列のアクセスコントロール」](#) を参照してください。

リードレプリカ

読み取り専用で使用されるデータベースのコピー。クエリをリードレプリカにルーティングして、プライマリデータベースへの負荷を軽減できます。

再設計

[「7 Rs」](#)を参照してください。

目標復旧時点 (RPO)

最後のデータリカバリポイントからの最大許容時間です。これにより、最後の回復時点からサービスが中断されるまでの間に許容できるデータ損失の程度が決まります。

目標復旧時間 (RTO)

サービスの中断から復旧までの最大許容遅延時間。

リファクタリング

[「7 Rs」](#)を参照してください。

リージョン

地理的エリア内の AWS リソースのコレクション。各 AWS リージョンは、耐障害性、安定性、耐障害性を提供するために、他のとは独立しています。詳細については、[AWS リージョン「アカウントで利用できるを指定する」](#)を参照してください。

回帰

数値を予測する機械学習手法。例えば、「この家はどれくらいの値段で売れるでしょうか?」という問題を解決するために、機械学習モデルは、線形回帰モデルを使用して、この家に関する既知の事実 (平方フィートなど) に基づいて家の販売価格を予測できます。

リホスト

[「7 Rs」](#)を参照してください。

リリース

デプロイプロセスで、変更を本番環境に昇格させること。

再配置

[「7 Rs」](#)を参照してください。

プラットフォーム変更

[「7 Rs」](#)を参照してください。

再購入

[「7 Rs」](#)を参照してください。

回復性

中断に抵抗または回復するアプリケーションの機能。[高可用性](#)と[ディザスタリカバリ](#)は、で回復性を計画する際の一般的な考慮事項です AWS クラウド。詳細については、[AWS クラウド「レジリエンス」](#)を参照してください。

リソースベースのポリシー

Amazon S3 バケット、エンドポイント、暗号化キーなどのリソースにアタッチされたポリシー。このタイプのポリシーは、アクセスが許可されているプリンシパル、サポートされているアクション、その他の満たすべき条件を指定します。

実行責任者、説明責任者、協業先、報告先 (RACI) に基づくマトリックス

移行活動とクラウド運用に関わるすべての関係者の役割と責任を定義したマトリックス。マトリックスの名前は、マトリックスで定義されている責任の種類、すなわち責任 (R)、説明責任 (A)、協議 (C)、情報提供 (I) に由来します。サポート (S) タイプはオプションです。サポートを含めると、そのマトリックスは RASCI マトリックスと呼ばれ、サポートを除外すると RACI マトリックスと呼ばれます。

レスポンスコントロール

有害事象やセキュリティベースラインからの逸脱について、修復を促すように設計されたセキュリティコントロール。詳細については、Implementing security controls on AWSの[Responsive controls](#)を参照してください。

保持

[「7 Rs」](#)を参照してください。

廃止

[「7 Rs」](#)を参照してください。

取得拡張生成 (RAG)

[LLM](#) がレスポンスを生成する前にトレーニングデータソースの外部にある信頼できるデータソースを参照する[生成 AI](#) テクノロジー。たとえば、RAG モデルは、組織のナレッジベースまたはカスタムデータのセマンティック検索を実行する場合があります。詳細については、[「RAG とは」](#)を参照してください。

ローテーション

攻撃者が認証情報にアクセスすることをより困難にするために、[シークレット](#)を定期的に更新するプロセス。

行と列のアクセス制御 (RCAC)

アクセスルールが定義された、基本的で柔軟な SQL 表現の使用。RCAC は行権限と列マスクで構成されています。

RPO

[「目標復旧時点」](#)を参照してください。

RTO

[目標復旧時間](#)を参照してください。

ランブック

特定のタスクを実行するために必要な手動または自動化された一連の手順。これらは通常、エラー率の高い反復操作や手順を合理化するために構築されています。

S

SAML 2.0

多くの ID プロバイダー (IdP) が使用しているオープンスタンダード。この機能を使用すると、フェデレーテッドシングルサインオン (SSO) が有効になるため、ユーザーは組織内のすべてのユーザーを IAM で作成しなくても、AWS Management Console にログインしたり AWS 、 API オペレーションを呼び出すことができます。SAML 2.0 ベースのフェデレーションの詳細については、IAM ドキュメントの[SAML 2.0 ベースのフェデレーションについて](#)を参照してください。

SCADA

[「監視コントロールとデータ取得」](#)を参照してください。

SCP

[「サービスコントロールポリシー」](#)を参照してください。

シークレット

暗号化された形式で保存する AWS Secrets Manager パスワードやユーザー認証情報などの機密情報または制限付き情報。シークレット値とそのメタデータで構成されます。シークレット値は、バイナリ、1 つの文字列、または複数の文字列にすることができます。詳細については、[Secrets Manager ドキュメントの「Secrets Manager シークレットの内容」](#)を参照してください。

設計によるセキュリティ

開発プロセス全体でセキュリティを考慮するシステムエンジニアリングアプローチ。

セキュリティコントロール

脅威アクターによるセキュリティ脆弱性の悪用を防止、検出、軽減するための、技術上または管理上のガードレール。セキュリティコントロールには、[予防的](#)、[検出的](#)、[応答的](#)、[プロ](#)アクティブの 4 つの主なタイプがあります。

セキュリティ強化

アタックサーフェスを狭めて攻撃への耐性を高めるプロセス。このプロセスには、不要になったリソースの削除、最小特権を付与するセキュリティのベストプラクティスの実装、設定ファイル内の不要な機能の無効化、といったアクションが含まれています。

Security Information and Event Management (SIEM) システム

セキュリティ情報管理 (SIM) とセキュリティイベント管理 (SEM) のシステムを組み合わせたツールとサービス。SIEM システムは、サーバー、ネットワーク、デバイス、その他ソースからデータを収集、モニタリング、分析して、脅威やセキュリティ違反を検出し、アラートを発信します。

セキュリティレスポンスの自動化

セキュリティイベントに自動的に応答または修復するように設計された、事前定義されたプログラムされたアクション。これらの自動化は、セキュリティのベストプラクティスを実装するのに役立つ[検出的](#)または[応答的](#)な AWS セキュリティコントロールとして機能します。自動応答アクションの例としては、VPC セキュリティグループの変更、Amazon EC2 インスタンスへのパッチ適用、認証情報の更新などがあります。

サーバー側の暗号化

送信先にあるデータの、それ AWS のサービス を受け取る による暗号化。

サービスコントロールポリシー (SCP)

AWS Organizationsの組織内の、すべてのアカウントのアクセス許可を一元的に管理するポリシー。SCP は、管理者がユーザーまたはロールに委任するアクションに、ガードレールを定義したり、アクションの制限を設定したりします。SCP は、許可リストまたは拒否リストとして、許可または禁止するサービスやアクションを指定する際に使用できます。詳細については、AWS Organizations ドキュメントの「[サービスコントロールポリシー](#)」を参照してください。

サービスエンドポイント

のエントリポイントの URL AWS のサービス。ターゲットサービスにプログラムで接続するには、エンドポイントを使用します。詳細については、AWS 全般のリファレンスの「[AWS のサービス エンドポイント](#)」を参照してください。

サービスレベルアグリーメント (SLA)

サービスのアップタイムやパフォーマンスなど、IT チームがお客様に提供すると約束したものを明示した合意書。

サービスレベルインジケータ (SLI)

エラー率、可用性、スループットなど、サービスのパフォーマンス側面の測定。

サービスレベルの目標 (SLO)

サービス [レベルのインジケータ](#) によって測定される、サービスの状態を表すターゲットメトリクス。

責任共有モデル

クラウドのセキュリティとコンプライアンス AWS について と共有する責任を説明するモデル。AWS はクラウドのセキュリティを担当しますが、お客様はクラウドのセキュリティを担当します。詳細については、[責任共有モデル](#) を参照してください。

SIEM

[セキュリティ情報とイベント管理システム](#) を参照してください。

単一障害点 (SPOF)

システムを中断する可能性のあるアプリケーションの 1 つの重要なコンポーネントの障害。

SLA

[「サービスレベルの契約」](#) を参照してください。

SLI

[「サービスレベルインジケータ」](#) を参照してください。

SLO

[「サービスレベルの目標」](#) を参照してください。

スプリットアンドシードモデル

モダナイゼーションプロジェクトのスケーリングと加速のためのパターン。新機能と製品リリースが定義されると、コアチームは解放されて新しい製品チームを作成します。これにより、お客様の組織の能力とサービスの拡張、デベロッパーの生産性の向上、迅速なイノベーションのサポートに役立ちます。詳細については、『』の [「アプリケーションをモダナイズするための段階的アプローチ AWS クラウド」](#) を参照してください。

SPOF

[単一障害点](#)を参照してください。

スタースキーマ

1 つの大きなファクトテーブルを使用してトランザクションデータまたは測定データを保存し、1 つ以上の小さなディメンションテーブルを使用してデータ属性を保存するデータベース組織構造。この構造は、[データウェアハウス](#)またはビジネスインテリジェンスの目的で使用するよう設計されています。

strangler fig パターン

レガシーシステムが廃止されるまで、システム機能を段階的に書き換えて置き換えることにより、モノリシックシステムをモダナイズするアプローチ。このパターンは、宿主の樹木から根を成長させ、最終的にその宿主を包み込み、宿主に取って代わるイチジクのつるを例えています。そのパターンは、モノリシックシステムを書き換えるときのリスクを管理する方法として [Martin Fowler により提唱されました](#)。このパターンの適用方法の例については、[コンテナと Amazon API Gateway を使用して、従来の Microsoft ASP.NET \(ASMX\) ウェブサービスを段階的にモダナイズ](#)を参照してください。

サブネット

VPC 内の IP アドレスの範囲。サブネットは、1 つのアベイラビリティゾーンに存在する必要があります。

監視制御とデータ収集 (SCADA)

製造では、ハードウェアとソフトウェアを使用して物理アセットと本番稼働をモニタリングするシステム。

対称暗号化

データの暗号化と復号に同じキーを使用する暗号化のアルゴリズム。

合成テスト

ユーザーとのやり取りをシミュレートして潜在的な問題を検出したり、パフォーマンスをモニタリングしたりする方法でシステムをテストします。[Amazon CloudWatch Synthetics](#) を使用して、これらのテストを作成できます。

システムプロンプト

[LLM](#) にコンテキスト、指示、またはガイドラインを提供して動作を指示する手法。システムプロンプトは、コンテキストを設定し、ユーザーとのやり取りのルールを確立するのに役立ちます。

T

tags

AWS リソースを整理するためのメタデータとして機能するキーと値のペア。タグは、リソースの管理、識別、整理、検索、フィルタリングに役立ちます。詳細については、「[AWS リソースのタグ付け](#)」を参照してください。

ターゲット変数

監督された機械学習でお客様が予測しようとしている値。これは、結果変数のことも指します。例えば、製造設定では、ターゲット変数が製品の欠陥である可能性があります。

タスクリスト

ランブックの進行状況を追跡するために使用されるツール。タスクリストには、ランブックの概要と完了する必要がある一般的なタスクのリストが含まれています。各一般的なタスクには、推定所要時間、所有者、進捗状況が含まれています。

テスト環境

[「環境」](#)を参照してください。

トレーニング

お客様の機械学習モデルに学習するデータを提供すること。トレーニングデータには正しい答えが含まれている必要があります。学習アルゴリズムは入力データ属性をターゲット (お客様が予測したい答え) にマッピングするトレーニングデータのパターンを検出します。これらのパターンをキャプチャする機械学習モデルを出力します。そして、お客様が機械学習モデルを使用して、ターゲットがわからない新しいデータでターゲットを予測できます。

トランジットゲートウェイ

VPC とオンプレミスネットワークを相互接続するために使用できる、ネットワークの中継ハブ。詳細については、AWS Transit Gateway ドキュメントの「[トランジットゲートウェイとは](#)」を参照してください。

トランクベースのワークフロー

デベロッパーが機能ブランチで機能をローカルにビルドしてテストし、その変更をメインブランチにマージするアプローチ。メインブランチはその後、開発環境、本番前環境、本番環境に合わせて順次構築されます。

信頼されたアクセス

ユーザーに代わって AWS Organizations およびそのアカウントで組織内でタスクを実行するために指定したサービスにアクセス許可を付与します。信頼されたサービスは、サービスにリンクされたロールを必要とときに各アカウントに作成し、ユーザーに代わって管理タスクを実行します。詳細については、ドキュメントの「[を他の AWS のサービス AWS Organizations で使用する AWS Organizations](#)」を参照してください。

チューニング

機械学習モデルの精度を向上させるために、お客様のトレーニングプロセスの側面を変更する。例えば、お客様が機械学習モデルをトレーニングするには、ラベル付けセットを生成し、ラベルを追加します。これらのステップを、異なる設定で複数回繰り返して、モデルを最適化します。

ツーピザチーム

2 枚のピザで養うことができるくらい小さな DevOps チーム。ツーピザチームの規模では、ソフトウェア開発におけるコラボレーションに最適な機会が確保されます。

U

不確実性

予測機械学習モデルの信頼性を損なう可能性がある、不正確、不完全、または未知の情報を指す概念。不確実性には、次の 2 つのタイプがあります。認識論的不確実性は、限られた、不完全なデータによって引き起こされ、弁論的不確実性は、データに固有のノイズとランダム性によって引き起こされます。詳細については、[深層学習システムにおける不確実性の定量化](#) ガイドを参照してください。

未分化なタスク

ヘビーリフティングとも呼ばれ、アプリケーションの作成と運用には必要だが、エンドユーザーに直接的な価値をもたらさなかったり、競争上の優位性をもたらしたりしない作業です。未分化なタスクの例としては、調達、メンテナンス、キャパシティプランニングなどがあります。

上位環境

[??? 「環境」](#) を参照してください。

V

バキューミング

ストレージを再利用してパフォーマンスを向上させるために、増分更新後にクリーンアップを行うデータベースのメンテナンス操作。

バージョンコントロール

リポジトリ内のソースコードへの変更など、変更を追跡するプロセスとツール。

VPC ピアリング

プライベート IP アドレスを使用してトラフィックをルーティングできる、2 つの VPC 間の接続。詳細については、Amazon VPC ドキュメントの「[VPC ピア機能とは](#)」を参照してください。

脆弱性

システムのセキュリティを脅かすソフトウェアまたはハードウェアの欠陥。

W

ウォームキャッシュ

頻繁にアクセスされる最新の関連データを含むバッファキャッシュ。データベースインスタンスはバッファキャッシュから、メインメモリまたはディスクからよりも短い時間で読み取りを行うことができます。

ウォームデータ

アクセス頻度の低いデータ。この種類のデータをクエリする場合、通常は適度に遅いクエリでも問題ありません。

ウィンドウ関数

現在のレコードに何らかの形で関連する行のグループに対して計算を実行する SQL 関数。ウィンドウ関数は、移動平均の計算や、現在の行の相対位置に基づく行の値へのアクセスなどのタスクの処理に役立ちます。

ワークロード

ビジネス価値をもたらすリソースとコード (顧客向けアプリケーションやバックエンドプロセスなど) の総称。

ワークストリーム

特定のタスクセットを担当する移行プロジェクト内の機能グループ。各ワークストリームは独立していますが、プロジェクト内の他のワークストリームをサポートしています。たとえば、ポートフォリオワークストリームは、アプリケーションの優先順位付け、ウェーブ計画、および移行メタデータの収集を担当します。ポートフォリオワークストリームは、これらの設備を移行ワークストリームで実現し、サーバーとアプリケーションを移行します。

WORM

「[Write Once](#)」、「[Read Many](#)」を参照してください。

WQF

[AWS「ワークロード認定フレームワーク」](#)を参照してください。

Write Once, Read Many (WORM)

データを 1 回書き込み、データの削除や変更を防ぐストレージモデル。承認されたユーザーは、必要な回数だけデータを読み取ることができますが、変更することはできません。このデータストレージインフラストラクチャは [イミュータブル](#) と見なされます。

Z

ゼロデイエクスプロイト

[ゼロデイ脆弱性](#) を利用する攻撃、通常はマルウェア。

ゼロデイ脆弱性

実稼働システムにおける未解決の欠陥または脆弱性。脅威アクターは、このような脆弱性を利用してシステムを攻撃する可能性があります。開発者は、よく攻撃の結果で脆弱性に気付きます。

ゼロショットプロンプト

[LLM](#) にタスクを実行する手順を提供しますが、タスクのガイドに役立つ例 (ショット) はありません。LLM は、事前トレーニング済みの知識を使用してタスクを処理する必要があります。ゼロショットプロンプトの有効性は、タスクの複雑さとプロンプトの品質によって異なります。「[数ショットプロンプト](#)」も参照してください。

ゾンビアプリケーション

平均 CPU およびメモリ使用率が 5% 未満のアプリケーション。移行プロジェクトでは、これらのアプリケーションを廃止するのが一般的です。

翻訳は機械翻訳により提供されています。提供された翻訳内容と英語版の間で齟齬、不一致または矛盾がある場合、英語版が優先します。