



調校 AWS Glue 和 Amazon EMR Spark 任務的 Spark SQL 查詢

AWS 方案指引



AWS 方案指引: 調校 AWS Glue 和 Amazon EMR Spark 任務的 Spark SQL 查詢

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

Amazon 的商標和商業外觀不得用於任何非 Amazon 的產品或服務，也不能以任何可能造成客戶混淆、任何貶低或使 Amazon 名譽受損的方式使用 Amazon 的商標和商業外觀。所有其他非 Amazon 擁有的商標均為其各自擁有者的財產，這些擁有者可能附屬於 Amazon，或與 Amazon 有合作關係，亦或受到 Amazon 贊助。

Table of Contents

簡介	1
目標對象	1
架構	2
最佳實務	5
使用聯結提示	5
BROADCAST	5
MERGE	6
SHUFFLE_HASH	8
SHUFFLE_REPLICATE_NL	10
使用 SQL 視窗函數	10
快取時使用單欄式格式	10
使用分割提示	11
使用 Adaptive Query Execution	11
刪除動態分割區	11
使用成本型最佳化工具	14
使用單欄格式	14
常見問答集	15
使用哪些最佳實務來擴展Apache Spark任務和分割區資料 AWS Glue ?	15
如何在 Amazon EMR 或 上提交Spark應用程式時解決可能的out-of-memory錯誤，並防止這些錯誤 AWS Glue ?	15
如何管理 中的記憶體和 CPU AWS Glue ?	15
資源	16
AWS 文件	16
Spark 文件	16
文件歷史紀錄	17
詞彙表	18
#	18
A	18
B	21
C	22
D	25
E	28
F	30
G	31

H	32
I	33
L	35
M	36
O	40
P	42
Q	44
R	45
S	47
T	50
U	51
V	52
W	52
Z	53
.....	liv

調校 AWS Glue 和 Amazon EMR Spark 任務的 Spark SQL 查詢

Phani Alapaty 和 Ravikiran Rao , Amazon Web Services (AWS)

2024 年 1 月 ([文件歷史記錄](#))

[Spark SQL](#) 是處理結構化資料的 Apache Spark 模組。[Amazon EMR](#) 和 [AWS Glue](#) 任務使用 Spark SQL 來處理、轉換和載入資料。與基本 Spark 彈性分散式資料集 (RDD) API 不同，Spark SQL 這些介面可提供 Spark 更多有關資料結構和執行中運算的資訊給。在內部，Spark SQL 會使用此額外資訊來執行其他查詢最佳化。有幾種方式可以與 互動 Spark SQL，包括 SQL 和資料集 API。

聯結資料是擷取、轉換或將資料載入物件存放區或資料庫時，您可以執行的最常見和重要操作之一。加入時，您需要考慮效能。當部分聯結、分析或彙總操作用盡記憶體時，有多個案例，例如大型網路傳輸。這可能會導致任務 AWS Glue Spark 失敗。

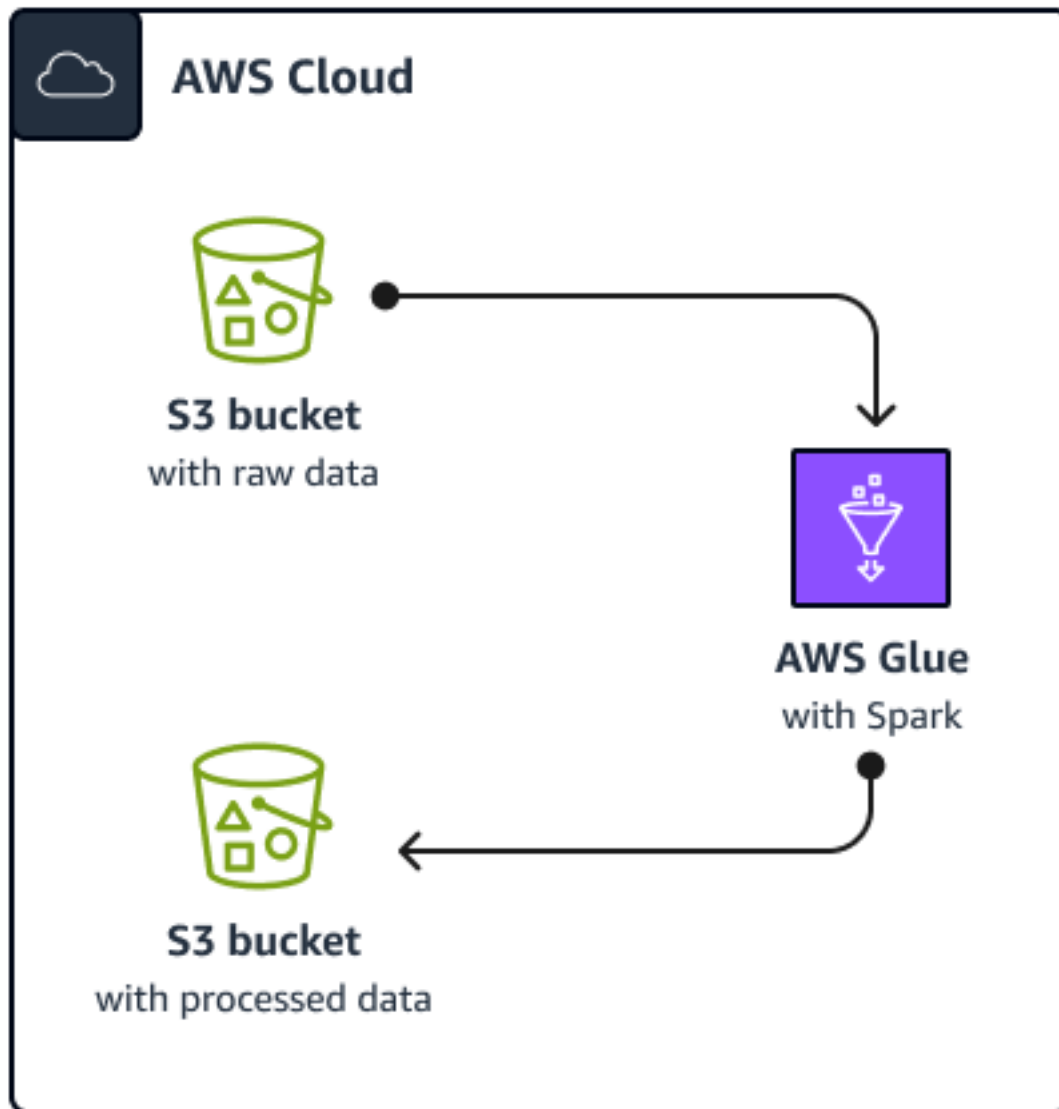
本指南提供的最佳實務可協助您調校 AWS Glue 或 Amazon EMR 任務的 Spark SQL 聯結查詢。Spark 提供許多組態選項，可改善 Spark SQL 工作負載的效能。這些調整可以透過程式設計方式完成，或者您可以使用 `spark-submit` 命令在全域層級套用。本指南說明其中一些組態，讓您可以改善或微調 Spark SQL 查詢和應用程式的效能。本指南中的建議是以 Professional Services 用來改善 Spark SQL 查詢和應用程式效能的組態 AWS 為基礎。

目標對象

本指南可協助架構師、資料工程師、資料科學家和開發人員了解可改善 Spark SQL 查詢效能的 Spark SQL 組態選項。

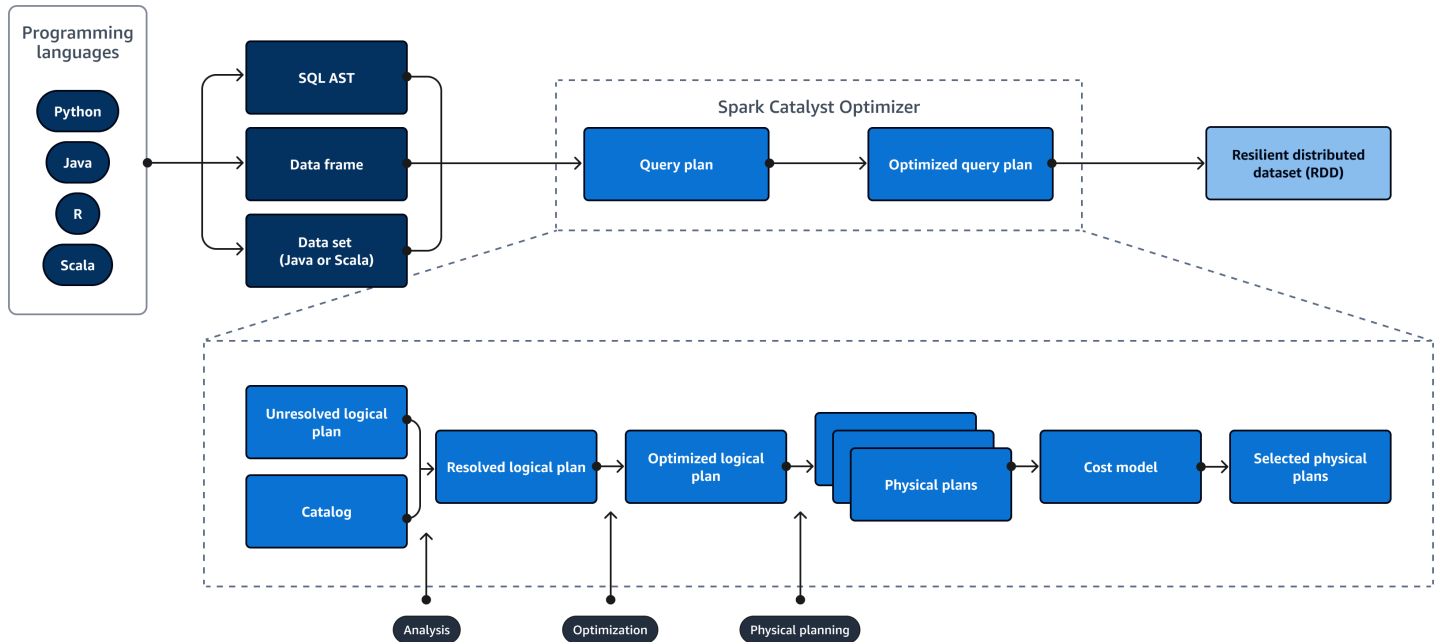
架構

下圖說明本指南中所述解決方案的架構。AWS Glue 任務會從 [Amazon Simple Storage Service \(Amazon S3\)](#) 儲存貯體讀取資料，這是雲端型物件儲存服務，可協助您儲存、保護和擷取資料。您可以透過 AWS Management Console、[AWS Command Line Interface \(AWS CLI\)](#) 或 AWS Glue API 啟動任務 AWS Glue Spark SQL。任務 AWS Glue Spark SQL 會在 Amazon S3 儲存貯體中處理原始資料，然後將處理的資料存放在不同的儲存貯體中。



舉例來說，本指南說明以 Python 和 Spark SQL() 撰寫的基本 [AWS Glue Spark SQL 任務](#) PySpark。此 AWS Glue 任務用於示範 Spark SQL 調校的最佳實務。雖然本指南著重於 AWS Glue，但本指南中的最佳實務也適用於 [Amazon EMR Spark SQL 任務](#)。

下圖說明 Spark SQL 查詢的生命週期。Spark SQL Catalyst Optimizer 會產生查詢計畫。查詢計畫是一系列步驟，如指示，用於存取 SQL 關聯式資料庫系統中的資料。若要開發效能最佳化 Spark SQL 的查詢計畫，第一步是檢視 EXPLAIN 計畫、解譯計畫，然後調整計畫。您可以使用 Spark SQL 使用者介面 (UI) 或 Spark SQL 歷史記錄伺服器來視覺化計畫。



Spark Catalyst Optimizer 會將初始查詢計畫轉換為最佳化的查詢計畫，如下所示：

1. 分析和宣告式 APIs – 分析階段是第一步。未解決的邏輯計畫，其中在 SQL 查詢中參考的物件不明或不符合輸入資料表，會產生未繫結屬性和資料類型。Spark SQL Catalyst Optimizer 接著會套用一組規則來建置邏輯計畫。SQL 剖析器可以產生 SQL 抽象語法樹 (AST)，並提供此做為邏輯計畫的輸入。輸入也可能是使用 API 建構的資料框架或資料集物件。下表顯示何時應使用 SQL、資料框架或資料集。

	SQL	資料框架	資料集
語法錯誤	執行期	編譯時間	編譯時間
分析錯誤	執行期	執行期	編譯時間

如需輸入類型的詳細資訊，請檢閱以下內容：

- 資料集 API 提供類型版本。這會因為高度依賴使用者定義的 lambda 函數而降低效能。RDD 或資料集會以靜態方式輸入。例如，當您定義 RDD 時，您需要明確提供結構描述定義。

- 資料框架 API 提供未輸入的關聯式操作。動態輸入資料影格。類似於 RDD，當您定義資料影格時，結構描述會保持不變。資料仍然是結構化的。不過，此資訊僅在執行時間提供。這可讓編譯器即時撰寫類似 SQL 的陳述式並定義新的資料欄。例如，它可以將資料欄附加到現有的資料框架，而不需要為每個操作定義新的類別。
 - Spark SQL 查詢會在執行時間評估語法和分析錯誤，以提供更快的執行時間。
2. 目錄 – Spark SQL 使用 Apache Hive Metastore (HMS) 管理持久性關聯式實體的中繼資料，例如資料庫、資料表、資料欄和分割區。
 3. 最佳化 – 最佳化工具會使用啟發式和成本來重寫查詢計劃。它會執行下列動作來產生最佳化的邏輯計畫：
 - a. 刪除資料欄
 - b. 下推述詞
 - c. 重新排序聯結
 4. 實體計畫和規劃器 – Spark SQL Catalyst Optimizer 會將邏輯計畫轉換為一組實體計畫。這表示它會將內容轉換為方法。
 5. 選取的實體計畫 – Spark SQL Catalyst Optimizer 會選取最具成本效益的實體計畫。
 6. 最佳化查詢計畫 – Spark SQL 執行效能最佳化和成本最佳化的查詢計畫。Spark SQL 記憶體管理會追蹤記憶體用量，並在任務和運算子之間分配記憶體。Spark SQL 編引擎可大幅改善 Spark SQL 應用程式的記憶體和 CPU 效率。它也會實作二進位資料模型處理，並直接在二進位資料上運作。這可避免還原序列化的需求，並大幅降低與資料轉換和還原序列化相關的額外負荷。

最佳實務

本節包含下列針對 Amazon EMR 和 調校 Spark SQL 查詢的最佳實務 AWS Glue：

- [在 中 使用聯結提示 Spark SQL](#)
- [使用 SQL 視窗函數而非 join 和 groupBy](#)
- [快取時使用單欄式格式](#)
- [在 Spark3.0.0 中使用分割提示](#)
- [使用 Adaptive Query Execution](#)
- [刪除動態分割區](#)
- [使用 Spark 成本型最佳化工具](#)
- [使用單欄式格式以獲得更好的查詢效能](#)

在 中 使用聯結提示 Spark SQL

使用 Spark3.0，您可以指定 Spark 要在執行時間使用的聯結演算法類型。聯結策略提示、SHUFFLE_HASH、BROADCAST MERGE 和 SHUFFLE_REPLICATE_NL 會指示在與其他關係聯結時 Spark，在每個指定的關係上使用提示策略。本節詳細討論聯結策略提示。

BROADCAST

在廣播雜湊聯結中，其中一個資料集明顯小於另一個資料集。由於較小的資料集可以容納在記憶體中，因此會廣播到叢集中的所有執行器。廣播資料後，會執行標準雜湊聯結。廣播雜湊聯結分為兩個步驟：

- 廣播 – 較小的資料集會廣播到叢集中的所有執行器。
- 雜湊聯結 – 較小的資料集會在所有執行器間雜湊，然後使用較大的資料集聯結。

沒有 sort 或 merge 操作。將具有較小維度資料表的大型事實資料表聯結以執行星狀結構描述聯結時，廣播雜湊是最快的聯結演算法。下列範例示範廣播雜湊聯結的運作方式。無論 `spark.sql.autoBroadcastJoinThreshold` 屬性中指定的大小限制為何，都會廣播具有提示的聯結端。如果聯結的兩端都有廣播提示，則會廣播具有較小大小（根據統計資料）的提示。`spark.sql.autoBroadcastJoinThreshold` 屬性的預設值為 10 MB。這會為在執行聯結時廣播至所有工作者節點的資料表設定大小上限，以位元組為單位。

下列範例提供查詢、實體EXPLAIN計畫，以及查詢執行所需的時間。如果您使用BROADCASTJOIN提示來強制廣播聯結，查詢需要較少的處理時間，如第二個範例EXPLAIN計劃所示。

```
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
```

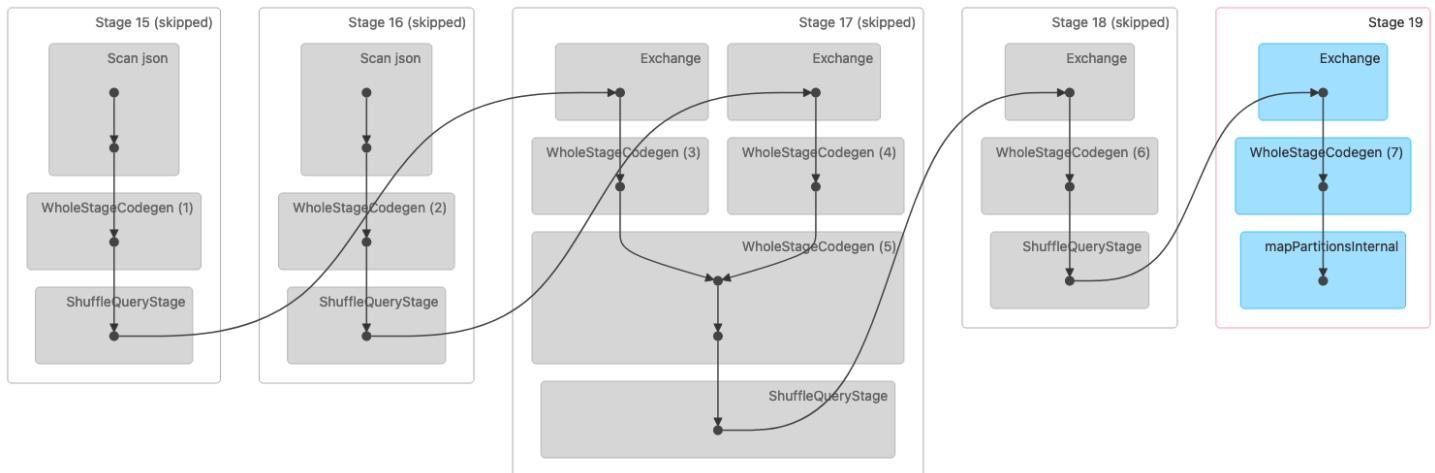
MERGE

當兩個資料集都很大且無法放入記憶體時，偏好使用隨機排序合併聯結。如名稱所示，此聯結包含下列三個階段：

- 隨機播放階段 – 聯結查詢中的兩個資料集都會隨機播放。
- 排序階段 – 記錄會依兩側的聯結索引鍵排序。
- 合併階段 – 根據聯結金鑰，反覆運算聯結條件的兩側。

下圖顯示隨機排序合併聯結的定向無環圖 (DAG) 視覺化。這兩個資料表會在前兩個階段讀取。在下一個階段（第 17 階段）中，它們會隨機播放、排序，然後在最後合併在一起。

注意：此影像中的某些階段顯示為略過，因為這些步驟已在先前的階段完成。該資料已快取或保留以供這些階段使用。



以下是指出排序合併聯結的實體計劃。

```

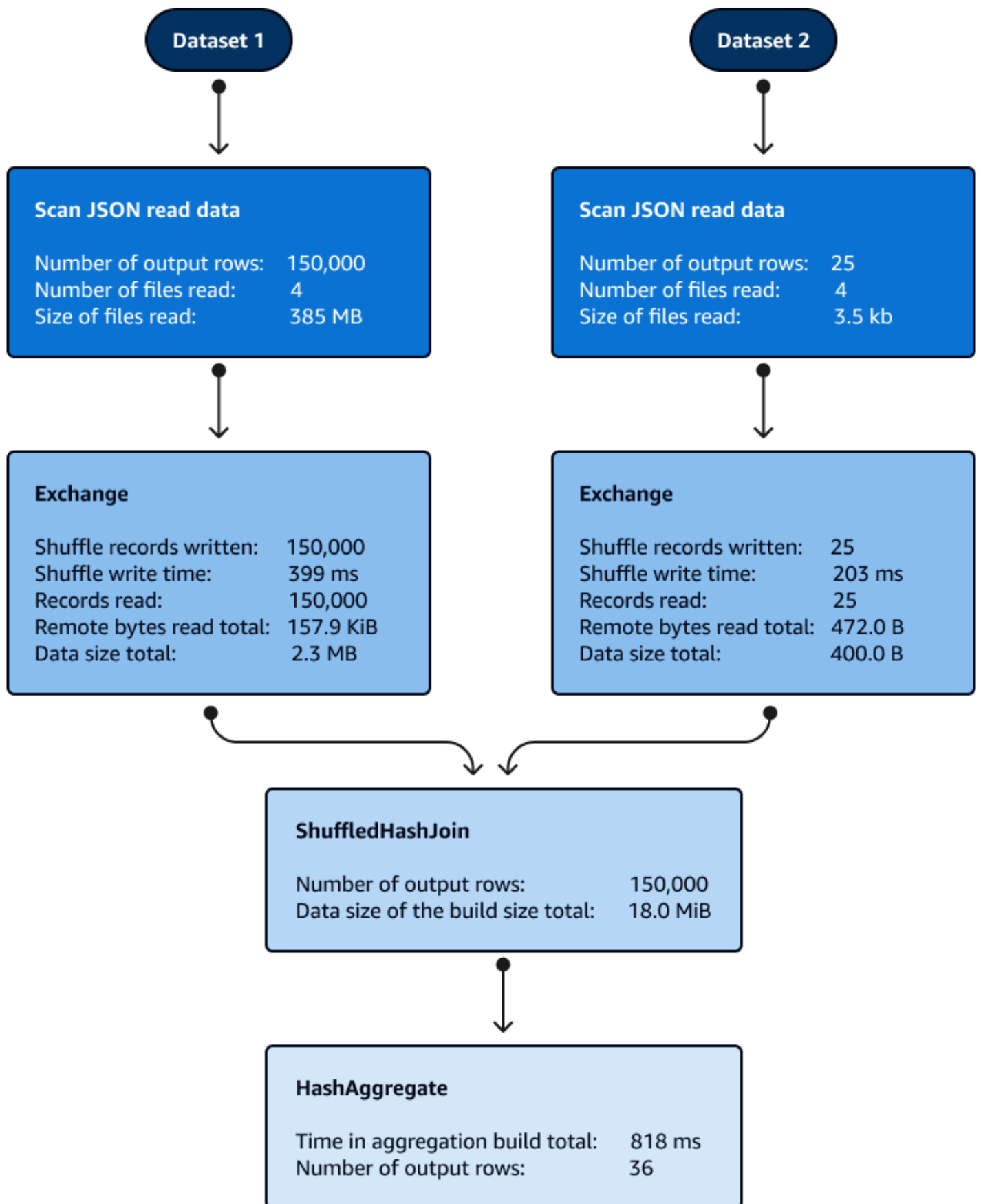
== Physical Plan ==
AdaptiveSparkPlan isFinalPlan=false
  +- SortMergeJoin [id#320L], [id#335L], Inner
    :- Sort [id#320L ASC NULLS FIRST], false, 0
    :   +- Exchange hashpartitioning(id#320L, 36), ENSURE_REQUIREMENTS, [id=#1018]
    :     +- Filter isnotnull(id#320L)
    :       +- Scan ExistingRDD[id#320L,col#321]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

如名稱所示，隨機播放雜湊聯結的運作方式是隨機播放這兩個資料集。來自兩側的相同金鑰最終會位於相同的分割區或任務中。資料隨機播放後，兩個資料集中最小的資料集會雜湊到儲存貯體中，然後在分割區中執行雜湊聯結。隨機雜湊聯結與廣播雜湊聯結不同，因為整個資料集不會廣播。隨機雜湊聯結分為兩個階段：

- 隨機播放階段 – 兩個資料集都會隨機播放。
- 雜湊聯結階段 – 資料較小的一端會經過雜湊處理、儲存，然後在所有分割區中加入較大一端的雜湊。

在分割區內加入隨機雜湊時，不需要排序。下圖顯示隨機雜湊聯結的階段。資料一開始會讀取，然後隨機播放，然後建立雜湊並用於聯結。



根據預設，當無法使用廣播雜湊聯結時，最佳化工具會選擇隨機雜湊聯結。根據廣播雜湊聯結閾值大小 (`spark.sql.autoBroadcastJoinThreshold`) 和所選隨機播放分割區的數量 (`spark.sql.shuffle.partitions`)，當邏輯 SQL 的單一分割區夠小而無法建置本機雜湊資料表時，它會使用隨機播放雜湊聯結。

SHUFFLE_REPLICATE_NL

Shuffle-and-Replicate 巢狀迴圈聯結也稱為笛卡爾產品聯結，運作方式非常類似廣播雜湊聯結，但資料集未廣播。

在此聯結演算法中，隨機播放不會參考真正的隨機播放，因為具有相同金鑰的記錄不會傳送至相同的分割區。反之，兩個資料集的整個分割區都會透過網路複製。當資料集中的分割區可用時，會執行巢狀迴圈聯結。如果第一個資料集中有記錄 X 數，且每個分割區中的第二個資料集中有記錄 Y 數，則第二個資料集中的每個記錄都會與第一個資料集中的每個記錄聯結。這會在每個分割區的循環 $X \times Y$ 時間繼續。

使用 SQL 視窗函數而非 join 和 groupBy

[視窗函數](#) (Spark 文件) 會在稱為視窗的一組資料列上執行計算，以某種方式與目前的記錄相關。例如，它們可以位於與目前資料列相同的分割區或影格中。視窗函數內建於 Spark 中，類似於彙總函數。彙總函數，例如 SUM 或 MAX，在一組資料列上操作，並計算每個群組的單一傳回值。視窗函數適用於處理任務，例如根據目前資料列的相對位置計算移動平均值或存取資料列的值。

Spark `groupBy` 函數會將資料收集到群組中，並對分組的資料執行彙總函數。當您使用 `groupBy`，會先 Spark 部分彙總資料，然後隨機播放減少的資料集。當您使用視窗函數時，整個資料集會隨機播放。視窗函數可以提供更快的執行時間。

在非常大型的資料集中，如果資料欄的基數很大，則建議使用視窗函數。不過，如果資料欄的基數很小，則資料彙總會很小，而且可以在聯結中廣播彙總結果。

快取時使用單欄式格式

Spark SQL 能夠以單欄式格式快取記憶體內資料表。

`spark.catalog.cacheTable("tableName")` 或 `dataFrame.cache()` 函數呼叫可用來快取記憶體內單欄式格式的資料表。然後，Spark SQL 引擎只會掃描必要的資料欄，並自動調整壓縮以減少記憶體和 CPU 使用量。您可以使用 `spark.catalog.uncacheTable("tableName")` 或從記憶體 `dataFrame.unpersist()` 中移除資料表。

在 Spark3.0.0 中使用分割提示

Spark [分割提示](#)可協助您調校效能並減少輸出檔案的數量。Spark SQL支援分割提示，例如 COALESCE、REPARTITION和 REPARTITION_BY_RANGE。這些提示類似於資料集 APIs，例如 coalesce、repartition和 repartitionByRange。下列提示可協助您控制 中的輸出檔案數量 Spark SQL，這可協助您調校效能：

- Coalesce - 將分割區數目減少為指定的分割區數目。分割區編號是COALESCE提示的唯一參數。
- 重新分割 - 使用指定的分割表達式重新分割至指定數量的分割區。REPARTITION 提示參數是分割區編號、資料欄名稱或兩者。
- 依範圍重新分割 - 使用指定的分割表達式重新分割至指定數量的分割區。資料欄名稱是REPARTITION_BY_RANGE提示的必要參數，分割區編號是選用的。
- 重新平衡 - 重新平衡查詢結果輸出分割區，讓每個分割區都是合理的大小。REBALANCE 提示參數是初始分割區編號、資料欄名稱，或兩者兼具。

使用 Adaptive Query Execution

Spark3.0 和更新版本包含稱為 的額外最佳化層。 [Adaptive Query Execution \(AQE\)](#)它會根據查詢執行時間期間收集的指標來最佳化查詢。它使用執行時間統計資料來挑選最有效率的執行計畫。此功能預設為在 3.2.0 版和更新Apache Spark版本中啟用。您可以將Spark組態設定為 來將其關閉spark.sql.adaptive.enabled=false。AQE 有三個主要功能：合併隨機播放後分割區、將排序合併聯結轉換為廣播聯結，以及最佳化扭曲聯結。

刪除動態分割區

Spark3.0 和更新版本包含動態分割區剔除 (DPP)。動態分割區剔除是 中的最佳化技術Spark，可防止讀取資料時掃描不必要的分割區。以下是有關 DPP 的一些重要須知：

- 它會檢查查詢篩選條件和述詞中請求的分割區值，並判斷需要哪些分割區才能滿足查詢。任何被視為不必要的分割區都會自動且透明地剔除。
- DPP 會略過不包含適用資料的分割區，以減少處理時間和資源使用率。這有助於僅Spark專注於相關分割區。
- 它適用於透過插入或增量載入新增的靜態分割區和動態產生的分割區。會Spark辨識新的分割區，並可繼續套用動態剔除。
- DPP 完全透明或不可見於開發人員。不需要特殊編碼即可啟用 DPP。它會在幕後自動發生，作為查詢計畫產生期間的最佳化。

以下是確保 DPP 有效運作的一些最佳實務：

- 在Spark資料框架操作的早期套用篩選條件，以使用述詞下推。這有助於提早使用分割區中繼資料 Spark來消除分割區。
- ANALYZE TABLE 頻繁執行以收集資料的統計資料。這可減少資料欄層級統計資料，Spark有助於更準確地判斷可以忽略哪些分割區。
- 避免過度分割您的資料。收集統計資料時，太多分割區可能會使驅動程式節點超載。目標為每個大型資料表 10-100 個分割區。
- 在聯結之前重新分割資料影格。這可防止需要移動所有資料的隨機聯結，並進一步最佳化讀取的資料量。
- 在要聯結的不同資料表中使用一致的分割區資料欄類型和命名。這有助於Spark更符合聯結最佳化的分割區。
- 使用 測試查詢EXPLAIN以確保 DPP 正在套用，並驗證是否需要新增調校。

在星星結構描述中，資料表分為兩種主要類型：事實資料表和維度資料表。維度資料表通常比事實資料表小得多。將事實資料表加入維度資料表時，DPP 會最佳化查詢計畫。它會從套用至維度資料表的任何篩選條件建立子查詢。它會廣播此子查詢，並從中建置雜湊資料表。然後，它會先將雜湊資料表套用至事實資料表的掃描階段，再讀取事實資料表資料。這有助於 DPP 減少必須從較大事實資料表讀取的資料量。

下列範例查詢顯示動作中的 DPP。查詢會從國家（印度）擷取訂單數量，並在事實資料表 (fact_orders) 和維度資料表 (nation) 之間包含內部聯結。fact_orders 資料表會依資料欄 分割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"
```

以下是EXPLAIN計劃中使用的步驟：

1. 掃描較小的維度資料表 (nation) 並依欄 篩選n_name = 'INDIA'。
2. 廣播上一個步驟的結果。
3. 建立篩選第一個步驟結果的子查詢。
4. 將其下推為 PartitionFilter使其僅掃描所需的事實資料表分割區，而不是完整資料表掃描。

以下是此 DPP 最佳化查詢的EXPLAIN計劃。

```

== 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 dynamicp...,
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>

```

即使資料 `o_nationkey` 欄上沒有新增直接篩選條件，由於 DPP 功能，Spark 會自動僅掃描需要的分割區，而不是整個資料表。

使用 Spark 成本型最佳化工具

Spark 2.09 版和更新版本可以使用成本型最佳化工具 (CBO)。CBO 會根據各種資料表統計資料，為查詢選取最便宜的執行計畫。CBO 會嘗試最佳化與 CPU 使用率和 I/O 相關的查詢執行，因此會盡快傳回。不過，在大多數情況下，資料統計資料通常不存在，特別是當統計資料收集比資料處理更昂貴時。即使統計資料可用，它們仍可能已過期。以下是使用 CBO 的最佳實務：

- 收集用於聯結、篩選條件和依分組等操作的所有資料欄的統計資料。這有助於最佳化工具做出更好的決策。使用 `ANALYZE TABLE` 收集統計資料。
- 如果資料快速變更，請提高統計資料收集的頻率。過時的統計資料可能會導致未最佳化的計畫。
- 從預設 CBO 參數開始，並僅在需要時進行調校。積極最佳化可以增加執行時間。
- 透過強制執行最佳化計畫並比較效能，確認 CBO 是否選擇最佳計畫。啟用最佳化工具記錄以比較計畫。
- 為了獲得更好的最佳化，請使用 CBO 搭配分割、儲存貯體和資料扭曲提示。分割可協助 CBO 擴展。
- 若要及早偵測計畫迴歸，請在啟用 CBO 後持續監控查詢效能。

使用單欄式格式以獲得更好的查詢效能

Spark 可以使用各種輸入檔案格式，例如 Apache Parquet、Optimized Row Columnar (ORC) 和 CSV。不過，在中 Parquet 效果最佳 Spark SQL。它提供更快的執行時間、更高的掃描輸送量、更低的磁碟 I/O 和更低的操作成本。Spark 可以透過按下推篩選條件使用 Parquet 檔案統計資料來自動篩選無用資料，例如最小/最大統計資料。另一方面，您可以啟用 parquet Spark 向量化讀取器，以批次讀取 Parquet 檔案。當您使用 Spark SQL 處理資料時，我們建議您盡可能使用 Parquet 檔案格式。

常見問答集

使用哪些最佳實務來擴展Apache Spark任務和分割區資料 AWS Glue ？

如需詳細資訊，請參閱 AWS 大數據部落格中的[使用 擴展Apache Spark任務和分割區資料的最佳實務 AWS Glue](#)。

如何在 Amazon EMR 或 上提交Spark應用程式時解決可能的out-of-memory錯誤，並防止這些錯誤 AWS Glue ？

如需詳細資訊，請參閱 上的 [EMR 最佳實務指南](#)GitHub。

如何管理 中的記憶體和 CPU AWS Glue ？

如需詳細資訊，請參閱 AWS 大數據部落格中的[最佳化 中的記憶體管理 AWS Glue](#)。

資源

AWS 文件

- [Apache Spark 任務效能調校 AWS Glue 的最佳實務](#)
- [AWS Command Line Interface \(AWS CLI\) 服務文件](#)
- [AWS Glue 服務文件](#)
- [Amazon Simple Storage Service \(Amazon S3\) 服務文件](#)
- [Amazon EMR 服務文件](#)

Spark 文件

- [效能調校](#)
- [SQL 查詢的聯結策略提示](#)

文件歷史紀錄

下表描述了本指南的重大變更。如果您想收到有關未來更新的通知，可以訂閱 [RSS 摘要](#)。

變更	描述	日期
初次出版	—	2024 年 1 月 4 日

AWS 規範性指引詞彙表

以下是 AWS Prescriptive Guidance 提供的策略、指南和模式中常用的術語。若要建議項目，請使用詞彙表末尾的提供意見回饋連結。

數字

7 R

將應用程式移至雲端的七種常見遷移策略。這些策略以 Gartner 在 2011 年確定的 5 R 為基礎，包括以下內容：

- 重構/重新架構 – 充分利用雲端原生功能來移動應用程式並修改其架構，以提高敏捷性、效能和可擴展性。這通常涉及移植作業系統和資料庫。範例：將您的現場部署 Oracle 資料庫遷移至 Amazon Aurora PostgreSQL 相容版本。
- 平台轉換 (隨即重塑) – 將應用程式移至雲端，並引入一定程度的優化以利用雲端功能。範例：將您的現場部署 Oracle 資料庫遷移至 中的 Amazon Relational Database Service (Amazon RDS) for Oracle AWS 雲端。
- 重新購買 (捨棄再購買) – 切換至不同的產品，通常從傳統授權移至 SaaS 模型。範例：將您的客戶關係管理 (CRM) 系統遷移至 Salesforce.com。
- 主機轉換 (隨即轉移) – 將應用程式移至雲端，而不進行任何變更以利用雲端功能。範例：將您的現場部署 Oracle 資料庫遷移至 中 EC2 執行個體上的 Oracle AWS 雲端。
- 重新放置 (虛擬機器監視器等級隨即轉移) – 將基礎設施移至雲端，無需購買新硬體、重寫應用程式或修改現有操作。您可以將伺服器從內部部署平台遷移到相同平台的雲端服務。範例：將 Microsoft Hyper-V 應用程式遷移至 AWS。
- 保留 (重新檢視) – 將應用程式保留在來源環境中。其中可能包括需要重要重構的應用程式，且您希望將該工作延遲到以後，以及您想要保留的舊版應用程式，因為沒有業務理由來進行遷移。
- 淘汰 – 解除委任或移除來源環境中不再需要的應用程式。

A

ABAC

請參閱[屬性型存取控制](#)。

抽象服務

請參閱 [受管服務](#)。

ACID

請參閱 [原子性、一致性、隔離性、持久性](#)。

主動-主動式遷移

一種資料庫遷移方法，其中來源和目標資料庫保持同步 (透過使用雙向複寫工具或雙重寫入操作)，且兩個資料庫都在遷移期間處理來自連接應用程式的交易。此方法支援小型、受控制批次的遷移，而不需要一次性切換。它更靈活，但比 [主動-被動遷移](#) 需要更多的工作。

主動-被動式遷移

一種資料庫遷移方法，其中來源和目標資料庫保持同步，但只有來源資料庫處理來自連接應用程式的交易，同時將資料複寫至目標資料庫。目標資料庫在遷移期間不接受任何交易。

彙總函數

在一組資料列上運作的 SQL 函數，會計算群組的單一傳回值。彙總函數的範例包括 SUM 和 MAX。

AI

請參閱 [人工智慧](#)。

AI Ops

請參閱 [人工智慧操作](#)。

匿名化

在資料集中永久刪除個人資訊的程序。匿名化有助於保護個人隱私權。匿名資料不再被視為個人資料。

反模式

經常用於重複性問題的解決方案，其中解決方案具有反生產力、無效或比替代解決方案更有效。

應用程式控制

一種安全方法，僅允許使用核准的應用程式，以協助保護系統免受惡意軟體攻擊。

應用程式組合

有關組織使用的每個應用程式的詳細資訊的集合，包括建置和維護應用程式的成本及其商業價值。此資訊是 [產品組合探索和分析程序](#) 的關鍵，有助於識別要遷移、現代化和優化的應用程式並排定其優先順序。

人工智慧 (AI)

電腦科學領域，致力於使用運算技術來執行通常與人類相關的認知功能，例如學習、解決問題和識別模式。如需詳細資訊，請參閱[什麼是人工智慧？](#)

人工智慧操作 (AIOps)

使用機器學習技術解決操作問題、減少操作事件和人工干預以及提高服務品質的程序。如需有關如何在 AWS 遷移策略中使用 AIOps 的詳細資訊，請參閱[操作整合指南](#)。

非對稱加密

一種加密演算法，它使用一對金鑰：一個用於加密的公有金鑰和一個用於解密的私有金鑰。您可以共用公有金鑰，因為它不用於解密，但對私有金鑰存取應受到高度限制。

原子性、一致性、隔離性、耐久性 (ACID)

一組軟體屬性，即使在出現錯誤、電源故障或其他問題的情況下，也能確保資料庫的資料有效性和操作可靠性。

屬性型存取控制 (ABAC)

根據使用者屬性 (例如部門、工作職責和團隊名稱) 建立精細許可的實務。如需詳細資訊，請參閱《AWS Identity and Access Management (IAM) 文件》中的[ABAC for AWS](#)。

授權資料來源

您存放主要版本資料的位置，被視為最可靠的資訊來源。您可以將授權資料來源中的資料複製到其他位置，以處理或修改資料，例如匿名、修訂或假名化資料。

可用區域

中的不同位置 AWS 區域，可隔離其他可用區域中的故障，並提供相同區域中其他可用區域的低成本、低延遲網路連線。

AWS 雲端採用架構 (AWS CAF)

的指導方針和最佳實務架構 AWS，可協助組織制定高效且有效的計劃，以成功地移至雲端。AWS CAF 將指導方針組織到六個重點領域：業務、人員、治理、平台、安全和營運。業務、人員和控管層面著重於業務技能和程序；平台、安全和操作層面著重於技術技能和程序。例如，人員層面針對處理人力資源 (HR)、人員配備功能和人員管理的利害關係人。因此，AWS CAF 為人員開發、訓練和通訊提供指引，協助組織做好成功採用雲端的準備。如需詳細資訊，請參閱[AWS CAF 網站](#)和[AWS CAF 白皮書](#)。

AWS 工作負載資格架構 (AWS WQF)

一種工具，可評估資料庫遷移工作負載、建議遷移策略，並提供工作預估值。AWS WQF 隨附於 AWS Schema Conversion Tool (AWS SCT)。它會分析資料庫結構描述和程式碼物件、應用程式程式碼、相依性和效能特性，並提供評估報告。

B

錯誤的機器人

旨在中斷或傷害個人或組織的[機器人](#)。

BCP

請參閱[業務持續性規劃](#)。

行為圖

資源行為的統一互動式檢視，以及一段時間後的互動。您可以將行為圖與 Amazon Detective 搭配使用來檢查失敗的登入嘗試、可疑的 API 呼叫和類似動作。如需詳細資訊，請參閱偵測文件中的[行為圖中的資料](#)。

大端序系統

首先儲存最高有效位元組的系統。另請參閱 [Endianness](#)。

二進制分類

預測二進制結果的過程 (兩個可能的類別之一)。例如，ML 模型可能需要預測諸如「此電子郵件是否是垃圾郵件？」等問題 或「產品是書還是汽車？」

Bloom 篩選條件

一種機率性、記憶體高效的資料結構，用於測試元素是否為集的成員。

藍/綠部署

一種部署策略，您可以在其中建立兩個不同但相同的環境。您可以在一個環境（藍色）中執行目前的應用程式版本，並在另一個環境（綠色）中執行新的應用程式版本。此策略可協助您快速復原，並將影響降至最低。

機器人

透過網際網路執行自動化任務並模擬人類活動或互動的軟體應用程式。有些機器人有用或有益，例如在網際網路上為資訊編製索引的 Web 爬蟲程式。有些其他機器人稱為惡意機器人，旨在中斷或傷害個人或組織。

殭屍網路

受到[惡意軟體](#)感染且受單一方控制之[機器人的](#)網路，稱為機器人繼承器或機器人運算子。殭屍網路是擴展機器人及其影響的最佳已知機制。

分支

程式碼儲存庫包含的區域。儲存庫中建立的第一個分支是主要分支。您可以從現有分支建立新分支，然後在新分支中開發功能或修正錯誤。您建立用來建立功能的分支通常稱為功能分支。當準備好發佈功能時，可以將功能分支合併回主要分支。如需詳細資訊，請參閱[關於分支](#) (GitHub 文件)。

碎片存取

在特殊情況下，以及透過核准的程序，讓使用者能夠快速存取他們通常無權存取 AWS 帳戶 的 。如需詳細資訊，請參閱 Well-Architected 指南中的 AWS [實作打破玻璃程序](#) 指標。

棕地策略

環境中的現有基礎設施。對系統架構採用棕地策略時，可以根據目前系統和基礎設施的限制來設計架構。如果正在擴展現有基礎設施，則可能會混合棕地和[綠地](#)策略。

緩衝快取

儲存最常存取資料的記憶體區域。

業務能力

業務如何創造價值 (例如，銷售、客戶服務或營銷)。業務能力可驅動微服務架構和開發決策。如需詳細資訊，請參閱在 [AWS 上執行容器化微服務](#) 白皮書的 [圍繞業務能力進行組織](#) 部分。

業務連續性規劃 (BCP)

一種解決破壞性事件 (如大規模遷移) 對營運的潛在影響並使業務能夠快速恢復營運的計畫。

C

CAF

請參閱[AWS 雲端採用架構](#)。

Canary 部署

版本對最終使用者的緩慢和增量版本。當您有信心時，您可以部署新版本並完全取代目前的版本。

CCoE

請參閱 [Cloud Center of Excellence](#)。

CDC

請參閱[變更資料擷取](#)。

變更資料擷取 (CDC)

追蹤對資料來源 (例如資料庫表格) 的變更並記錄有關變更的中繼資料的程序。您可以將 CDC 用於各種用途，例如稽核或複寫目標系統中的變更以保持同步。

混沌工程

故意引入故障或破壞性事件，以測試系統的彈性。您可以使用 [AWS Fault Injection Service \(AWS FIS\)](#) 執行實驗，為您的 AWS 工作負載帶來壓力，並評估其回應。

CI/CD

請參閱[持續整合和持續交付](#)。

分類

有助於產生預測的分類程序。用於分類問題的 ML 模型可預測離散值。離散值永遠彼此不同。例如，模型可能需要評估影像中是否有汽車。

用戶端加密

在目標 AWS 服務 接收資料之前，在本機加密資料。

雲端卓越中心 (CCoE)

一個多學科團隊，可推動整個組織的雲端採用工作，包括開發雲端最佳實務、調動資源、制定遷移時間表以及領導組織進行大規模轉型。如需詳細資訊，請參閱 AWS 雲端 企業策略部落格上的 [CCoE 文章](#)。

雲端運算

通常用於遠端資料儲存和 IoT 裝置管理的雲端技術。雲端運算通常連接到[邊緣運算](#)技術。

雲端操作模型

在 IT 組織中，用於建置、成熟和最佳化一或多個雲端環境的操作模型。如需詳細資訊，請參閱[建置您的雲端操作模型](#)。

採用雲端階段

組織在遷移至 時通常會經歷的四個階段 AWS 雲端：

- 專案 – 執行一些與雲端相關的專案以進行概念驗證和學習用途
- 基礎 – 進行基礎投資以擴展雲端採用 (例如，建立登陸區域、定義 CCoE、建立營運模型)

- 遷移 – 遷移個別應用程式
- 重塑 – 優化產品和服務，並在雲端中創新

這些階段由 Stephen Orban 於部落格文章 [The Journey Toward Cloud-First 和 Enterprise Strategy 部落格上的採用階段](#) 中定義。AWS 雲端 如需有關它們如何與 AWS 遷移策略相關的詳細資訊，請參閱 [遷移整備指南](#)。

CMDB

請參閱 [組態管理資料庫](#)。

程式碼儲存庫

透過版本控制程序來儲存及更新原始程式碼和其他資產 (例如文件、範例和指令碼) 的位置。常見的雲端儲存庫包括 GitHub 或 Bitbucket Cloud。程式碼的每個版本都稱為分支。在微服務結構中，每個儲存庫都專用於單個功能。單一 CI/CD 管道可以使用多個儲存庫。

冷快取

一種緩衝快取，它是空的、未填充的，或者包含過時或不相關的資料。這會影響效能，因為資料庫執行個體必須從主記憶體或磁碟讀取，這比從緩衝快取讀取更慢。

冷資料

很少存取且通常是歷史資料的資料。查詢這類資料時，通常可接受慢查詢。將此資料移至效能較低且成本較低的儲存層或類別，可以降低成本。

電腦視覺 (CV)

使用機器學習從數位影像和影片等視覺化格式分析和擷取資訊的 [AI](#) 欄位。例如，Amazon SageMaker AI 提供 CV 的影像處理演算法。

組態偏離

對於工作負載，組態會從預期狀態變更。這可能會導致工作負載變得不合規，而且通常是漸進和無意的。

組態管理資料庫 (CMDB)

儲存和管理有關資料庫及其 IT 環境的資訊的儲存庫，同時包括硬體和軟體元件及其組態。您通常在遷移的產品組合探索和分析階段使用 CMDB 中的資料。

一致性套件

您可以組合的 AWS Config 規則和修補動作集合，以自訂您的合規和安全檢查。您可以使用 YAML 範本，將一致性套件部署為 AWS 帳戶 和 區域中或整個組織的單一實體。如需詳細資訊，請參閱 AWS Config 文件中的 [一致性套件](#)。

持續整合和持續交付 (CI/CD)

自動化軟體發程序的來源、建置、測試、暫存和生產階段的程序。CI/CD 通常被描述為管道。CI/CD 可協助您將程序自動化、提升生產力、改善程式碼品質以及加快交付速度。如需詳細資訊，請參閱[持續交付的優點](#)。CD 也可表示持續部署。如需詳細資訊，請參閱[持續交付與持續部署](#)。

CV

請參閱[電腦視覺](#)。

D

靜態資料

網路中靜止的資料，例如儲存中的資料。

資料分類

根據重要性和敏感性來識別和分類網路資料的程序。它是所有網路安全風險管理策略的關鍵組成部分，因為它可以協助您確定適當的資料保護和保留控制。資料分類是 AWS Well-Architected Framework 中安全支柱的元件。如需詳細資訊，請參閱[資料分類](#)。

資料偏離

生產資料與用於訓練 ML 模型的資料之間有意義的變化，或輸入資料隨時間有意義的變更。資料偏離可以降低 ML 模型預測的整體品質、準確性和公平性。

傳輸中的資料

在您的網路中主動移動的資料，例如在網路資源之間移動。

資料網格

架構架構，提供分散式、分散式資料擁有權與集中式管理。

資料最小化

僅收集和處理嚴格必要資料的原則。在 中實作資料最小化 AWS 雲端 可以降低隱私權風險、成本和分析碳足跡。

資料周邊

AWS 環境中的一組預防性防護機制，可協助確保只有信任的身分才能從預期的網路存取信任的資源。如需詳細資訊，請參閱[在上建置資料周邊 AWS](#)。

資料預先處理

將原始資料轉換成 ML 模型可輕鬆剖析的格式。預處理資料可能意味著移除某些欄或列，並解決遺失、不一致或重複的值。

資料來源

在整個生命週期中追蹤資料的原始伺服器 and 歷史記錄的程序，例如資料的產生、傳輸和儲存方式。

資料主體

正在收集和處理其資料的個人。

資料倉儲

支援商業智慧的資料管理系統，例如 分析。資料倉儲通常包含大量歷史資料，通常用於查詢和分析。

資料庫定義語言 (DDL)

用於建立或修改資料庫中資料表和物件之結構的陳述式或命令。

資料庫處理語言 (DML)

用於修改 (插入、更新和刪除) 資料庫中資訊的陳述式或命令。

DDL

請參閱[資料庫定義語言](#)。

深度整體

結合多個深度學習模型進行預測。可以使用深度整體來獲得更準確的預測或估計預測中的不確定性。

深度學習

一個機器學習子領域，它使用多層人工神經網路來識別感興趣的輸入資料與目標變數之間的對應關係。

深度防禦

這是一種資訊安全方法，其中一系列的安全機制和控制項會在整個電腦網路中精心分層，以保護網路和其中資料的機密性、完整性和可用性。當您在 上採用此策略時 AWS，您可以在 AWS Organizations 結構的不同層新增多個控制項，以協助保護資源。例如，defense-in-depth方法可能會結合多重重要素驗證、網路分割和加密。

委派的管理員

在 [中 AWS Organizations](#)，相容的服務可以註冊 AWS 成員帳戶，以管理組織的帳戶和管理該服務的許可。此帳戶稱為該服務的委派管理員。如需詳細資訊和相容服務清單，請參閱 [AWS Organizations 文件中的可搭配 AWS Organizations 運作的服務](#)。

部署

在目標環境中提供應用程式、新功能或程式碼修正的程序。部署涉及在程式碼庫中實作變更，然後在應用程式環境中建置和執行該程式碼庫。

開發環境

請參閱 [環境](#)。

偵測性控制

一種安全控制，用於在事件發生後偵測、記錄和提醒。這些控制是第二道防線，提醒您注意繞過現有預防性控制的安全事件。如需詳細資訊，請參閱在 [AWS 上實作安全控制中的偵測性控制](#)。

開發值串流映射 (DVSM)

一種程序，用於識別對軟體開發生命週期中的速度和品質造成負面影響的限制並排定優先順序。DVSM 延伸了原本專為精簡製造實務設計的價值串流映射程序。它著重於透過軟體開發程序建立和移動價值所需的步驟和團隊。

數位分身

真實世界系統的虛擬呈現，例如建築物、工廠、工業設備或生產線。數位分身支援預測性維護、遠端監控和生產最佳化。

維度資料表

在 [星星結構描述](#) 中，較小的資料表包含有關事實資料表中量化資料的資料屬性。維度資料表屬性通常是文字欄位或離散數字，其行為類似於文字。這些屬性通常用於查詢限制、篩選和結果集標記。

災難

防止工作負載或系統在其主要部署位置中實現其業務目標的事件。這些事件可能是自然災難、技術故障或人為動作的結果，例如意外設定錯誤或惡意軟體攻擊。

災難復原 (DR)

您用來將 [災難](#) 造成的停機時間和資料遺失降至最低的策略和程序。如需詳細資訊，請參閱 [AWS Well-Architected Framework 中的上工作負載災難復原 AWS：雲端中的復原](#)。

DML

請參閱[資料庫處理語言](#)。

領域驅動的設計

一種開發複雜軟體系統的方法，它會將其元件與每個元件所服務的不斷發展的領域或核心業務目標相關聯。Eric Evans 在其著作 *Domain-Driven Design: Tackling Complexity in the Heart of Software* (Boston: Addison-Wesley Professional, 2003) 中介紹了這一概念。如需有關如何將領域驅動的設計與 strangler fig 模式搭配使用的資訊，請參閱[使用容器和 Amazon API Gateway 逐步現代化舊版 Microsoft ASP.NET \(ASMX\) Web 服務](#)。

DR

請參閱[災難復原](#)。

偏離偵測

追蹤與基準組態的偏差。例如，您可以使用 AWS CloudFormation 來偵測系統資源中的偏離，也可以使用 AWS Control Tower 來[偵測登陸區域中可能影響控管要求合規性的變更](#)。<https://docs.aws.amazon.com/AWSCloudFormation/latest/UserGuide/using-cfn-stack-drift.html>

DVSM

請參閱[開發值串流映射](#)。

E

EDA

請參閱[探索性資料分析](#)。

EDI

請參閱[電子資料交換](#)。

邊緣運算

提升 IoT 網路邊緣智慧型裝置運算能力的技術。與[雲端運算](#)相比，邊緣運算可以減少通訊延遲並改善回應時間。

電子資料交換 (EDI)

在組織之間自動交換商業文件。如需詳細資訊，請參閱[什麼是電子資料交換](#)。

加密

一種運算程序，可將人類可讀取的純文字資料轉換為加密文字。

加密金鑰

由加密演算法產生的隨機位元的加密字串。金鑰長度可能有所不同，每個金鑰的設計都是不可預測且唯一的。

端序

位元組在電腦記憶體中的儲存順序。大端序系統首先儲存最高有效位元組。小端序系統首先儲存最低有效位元組。

端點

請參閱 [服務端點](#)。

端點服務

您可以在虛擬私有雲端 (VPC) 中託管以與其他使用者共用的服務。您可以使用 [建立端點服務](#)，AWS PrivateLink 並將許可授予其他 AWS 帳戶 或 AWS Identity and Access Management (IAM) 委託人。這些帳戶或主體可以透過建立介面 VPC 端點私下連接至您的端點服務。如需詳細資訊，請參閱 Amazon Virtual Private Cloud (Amazon VPC) 文件中的[建立端點服務](#)。

企業資源規劃 (ERP)

一種系統，可自動化和和管理企業的關鍵業務流程（例如會計、[MES](#) 和專案管理）。

信封加密

使用另一個加密金鑰對某個加密金鑰進行加密的程序。如需詳細資訊，請參閱 AWS Key Management Service (AWS KMS) 文件中的[信封加密](#)。

環境

執行中應用程式的執行個體。以下是雲端運算中常見的環境類型：

- 開發環境 – 執行中應用程式的執行個體，只有負責維護應用程式的核心團隊才能使用。開發環境用來測試變更，然後再將開發環境提升到較高的環境。此類型的環境有時稱為測試環境。
- 較低的環境 – 應用程式的所有開發環境，例如用於初始建置和測試的開發環境。
- 生產環境 – 最終使用者可以存取的執行中應用程式的執行個體。在 CI/CD 管道中，生產環境是最後一個部署環境。
- 較高的環境 – 核心開發團隊以外的使用者可存取的所有環境。這可能包括生產環境、生產前環境以及用於使用者接受度測試的環境。

epic

在敏捷方法中，有助於組織工作並排定工作優先順序的功能類別。epic 提供要求和實作任務的高層級描述。例如，AWS CAF 安全概念包括身分和存取管理、偵測控制、基礎設施安全、資料保護和事件回應。如需有關 AWS 遷移策略中的 Epic 的詳細資訊，請參閱[計畫實作指南](#)。

ERP

請參閱[企業資源規劃](#)。

探索性資料分析 (EDA)

分析資料集以了解其主要特性的過程。您收集或彙總資料，然後執行初步調查以尋找模式、偵測異常並檢查假設。透過計算摘要統計並建立資料可視化來執行 EDA。

F

事實資料表

[星狀結構描述](#)中的中央資料表。它存放有關業務操作的量化資料。一般而言，事實資料表包含兩種類型的資料欄：包含度量的資料，以及包含維度資料表外部索引鍵的資料欄。

快速失敗

一種使用頻繁和增量測試來縮短開發生命週期的理念。這是敏捷方法的關鍵部分。

故障隔離界限

在中 AWS 雲端，像是可用區域 AWS 區域、控制平面或資料平面等邊界會限制故障的影響，並有助於改善工作負載的彈性。如需詳細資訊，請參閱[AWS 故障隔離界限](#)。

功能分支

請參閱[分支](#)。

特徵

用來進行預測的輸入資料。例如，在製造環境中，特徵可能是定期從製造生產線擷取的影像。

功能重要性

特徵對於模型的預測有多重要。這通常表示為可以透過各種技術來計算的數值得分，例如 Shapley Additive Explanations (SHAP) 和積分梯度。如需詳細資訊，請參閱[機器學習模型可解釋性 AWS](#)。

特徵轉換

優化 ML 程序的資料，包括使用其他來源豐富資料、調整值、或從單一資料欄位擷取多組資訊。這可讓 ML 模型從資料中受益。例如，如果將「2021-05-27 00:15:37」日期劃分為「2021」、「五月」、「週四」和「15」，則可以協助學習演算法學習與不同資料元件相關聯的細微模式。

少量擷取提示

在要求 [LLM](#) 執行類似的任務之前，提供少量示範任務和所需輸出的範例。此技術是內容內學習的應用程式，其中模型會從內嵌在提示中的範例 (快照) 中學習。對於需要特定格式、推理或網域知識的任務，少量的提示非常有效。另請參閱[零鏡頭提示](#)。

FGAC

請參閱[精細存取控制](#)。

精細存取控制 (FGAC)

使用多個條件來允許或拒絕存取請求。

閃切遷移

一種資料庫遷移方法，透過[變更資料擷取](#)使用連續資料複寫，以盡可能在最短的時間內遷移資料，而不是使用分階段方法。目標是將停機時間降至最低。

FM

請參閱[基礎模型](#)。

基礎模型 (FM)

大型深度學習神經網路，已針對廣義和未標記資料的大量資料集進行訓練。FMs 能夠執行各種一般任務，例如了解語言、產生文字和影像，以及以自然語言交談。如需詳細資訊，請參閱[什麼是基礎模型](#)。

G

生成式 AI

已針對大量資料進行訓練的 [AI](#) 模型子集，可使用簡單的文字提示建立新的內容和成品，例如影像、影片、文字和音訊。如需詳細資訊，請參閱[什麼是生成式 AI](#)。

地理封鎖

請參閱[地理限制](#)。

地理限制 (地理封鎖)

Amazon CloudFront 中的選項，可防止特定國家/地區的使用者存取內容分發。您可以使用允許清單或封鎖清單來指定核准和禁止的國家/地區。如需詳細資訊，請參閱 CloudFront 文件中的[限制內容的地理分佈](#)。

Gitflow 工作流程

這是一種方法，其中較低和較高環境在原始碼儲存庫中使用不同分支。Gitflow 工作流程會被視為舊版，而以[幹線為基礎的工作流程](#)是現代、偏好的方法。

黃金影像

系統或軟體的快照，做為部署該系統或軟體新執行個體的範本。例如，在製造中，黃金映像可用於在多個裝置上佈建軟體，並有助於提高裝置製造操作的速度、可擴展性和生產力。

綠地策略

新環境中缺乏現有基礎設施。對系統架構採用綠地策略時，可以選擇所有新技術，而不會限制與現有基礎設施的相容性，也稱為[棕地](#)。如果正在擴展現有基礎設施，則可能會混合棕地和綠地策略。

防護機制

有助於跨組織單位 (OU) 來管控資源、政策和合規的高層級規則。預防性防護機制會強制執行政策，以確保符合合規標準。透過使用服務控制政策和 IAM 許可界限來將其實作。偵測性防護機制可偵測政策違規和合規問題，並產生提醒以便修正。它們是透過使用 AWS Config AWS Security Hub、Amazon GuardDuty、Amazon Inspector AWS Trusted Advisor和自訂 AWS Lambda 檢查來實作。

H

HA

請參閱[高可用性](#)。

異質資料庫遷移

將來源資料庫遷移至使用不同資料庫引擎的目標資料庫 (例如，Oracle 至 Amazon Aurora)。異質遷移通常是重新架構工作的一部分，而轉換結構描述可能是一項複雜任務。[AWS 提供有助於結構描述轉換的 AWS SCT](#)。

高可用性 (HA)

在遇到挑戰或災難時，工作負載能夠在不介入的情況下持續運作。HA 系統的設計目的是自動容錯移轉、持續提供高品質的效能，以及處理不同的負載和故障，並將效能影響降至最低。

歷史現代化

一種方法，用於現代化和升級操作技術 (OT) 系統，以更好地滿足製造業的需求。歷史資料是一種資料庫，用於從工廠中的各種來源收集和存放資料。

保留資料

從用於訓練機器學習模型的資料集中保留的部分歷史標記資料。您可以使用保留資料，透過比較模型預測與保留資料來評估模型效能。

異質資料庫遷移

將您的來源資料庫遷移至共用相同資料庫引擎的目標資料庫 (例如，Microsoft SQL Server 至 Amazon RDS for SQL Server)。同質遷移通常是主機轉換或平台轉換工作的一部分。您可以使用原生資料庫公用程式來遷移結構描述。

熱資料

經常存取的資料，例如即時資料或最近的轉譯資料。此資料通常需要高效能儲存層或類別，才能提供快速的查詢回應。

修補程序

緊急修正生產環境中的關鍵問題。由於其緊迫性，通常會在典型 DevOps 發行工作流程之外執行修補程式。

超級護理期間

在切換後，遷移團隊在雲端管理和監控遷移的應用程式以解決任何問題的時段。通常，此期間的長度為 1-4 天。在超級護理期間結束時，遷移團隊通常會將應用程式的責任轉移給雲端營運團隊。

|

IaC

將[基礎設施視為程式碼](#)。

身分型政策

連接至一或多個 IAM 主體的政策，可定義其在 AWS 雲端環境中的許可。

閒置應用程式

90 天期間 CPU 和記憶體平均使用率在 5% 至 20% 之間的應用程式。在遷移專案中，通常會淘汰這些應用程式或將其保留在內部部署。

IIoT

請參閱[工業物聯網](#)。

不可變的基礎設施

為生產工作負載部署新基礎設施的模型，而不是更新、修補或修改現有基礎設施。不可變基礎設施本質上比[可變基礎設施](#)更一致、可靠且可預測。如需詳細資訊，請參閱 AWS Well-Architected Framework [中的使用不可變基礎設施部署](#)最佳實務。

傳入 (輸入) VPC

在 AWS 多帳戶架構中，接受、檢查和路由來自應用程式外部之網路連線的 VPC。[AWS 安全參考架構](#)建議您使用傳入、傳出和檢查 VPC 來設定網路帳戶，以保護應用程式與更廣泛的網際網路之間的雙向介面。

增量遷移

一種切換策略，您可以在其中將應用程式分成小部分遷移，而不是執行單一、完整的切換。例如，您最初可能只將一些微服務或使用者移至新系統。確認所有項目都正常運作之後，您可以逐步移動其他微服務或使用者，直到可以解除委任舊式系統。此策略可降低與大型遷移關聯的風險。

工業 4.0

[Klaus 於](#) 2016 年引進的術語，透過連線能力、即時資料、自動化、分析和 AI/ML 的進展，指製造程序的現代化。

基礎設施

應用程式環境中包含的所有資源和資產。

基礎設施即程式碼 (IaC)

透過一組組態檔案來佈建和管理應用程式基礎設施的程序。IaC 旨在協助您集中管理基礎設施，標準化資源並快速擴展，以便新環境可重複、可靠且一致。

工業物聯網 (IIoT)

在製造業、能源、汽車、醫療保健、生命科學和農業等產業領域使用網際網路連線的感測器和裝置。如需詳細資訊，請參閱[建立工業物聯網 \(IIoT\) 數位轉型策略](#)。

檢查 VPC

在 AWS 多帳戶架構中，集中式 VPC 可管理 VPCs 之間（在相同或不同的 AWS 區域）、網際網路和內部部署網路之間的網路流量檢查。[AWS 安全參考架構](#)建議您使用傳入、傳出和檢查 VPC 來設定網路帳戶，以保護應用程式與更廣泛的網際網路之間的雙向介面。

物聯網 (IoT)

具有內嵌式感測器或處理器的相連實體物體網路，其透過網際網路或本地通訊網路與其他裝置和系統進行通訊。如需詳細資訊，請參閱[什麼是 IoT？](#)

可解釋性

機器學習模型的一個特徵，描述了人類能夠理解模型的預測如何依賴於其輸入的程度。如需詳細資訊，請參閱[的機器學習模型可解釋性 AWS](#)。

IoT

請參閱[物聯網](#)。

IT 資訊庫 (ITIL)

一組用於交付 IT 服務並使這些服務與業務需求保持一致的最佳實務。ITIL 為 ITSM 提供了基礎。

IT 服務管理 (ITSM)

與組織的設計、實作、管理和支援 IT 服務關聯的活動。如需有關將雲端操作與 ITSM 工具整合的資訊，請參閱[操作整合指南](#)。

ITIL

請參閱[IT 資訊庫](#)。

ITSM

請參閱[IT 服務管理](#)。

L

標籤型存取控制 (LBAC)

強制存取控制 (MAC) 的實作，其中使用者和資料本身都會獲得明確指派的安全標籤值。使用者安全標籤和資料安全標籤之間的交集會決定使用者可以看到哪些資料列和資料欄。

登陸區域

登陸區域是架構良好的多帳戶 AWS 環境，可擴展且安全。這是一個起點，您的組織可以從此起點快速啟動和部署工作負載與應用程式，並對其安全和基礎設施環境充滿信心。如需有關登陸區域的詳細資訊，請參閱[設定安全且可擴展的多帳戶 AWS 環境](#)。

大型語言模型 (LLM)

預先訓練大量資料的深度學習 [AI](#) 模型。LLM 可以執行多個任務，例如回答問題、摘要文件、將文字翻譯成其他語言，以及完成句子。如需詳細資訊，請參閱[什麼是 LLMs](#)。

大型遷移

遷移 300 部或更多伺服器。

LBAC

請參閱[標籤型存取控制](#)。

最低權限

授予執行任務所需之最低許可的安全最佳實務。如需詳細資訊，請參閱 IAM 文件中的[套用最低權限許可](#)。

隨即轉移

請參閱 [7 個 R](#)。

小端序系統

首先儲存最低有效位元組的系統。另請參閱 [Endianness](#)。

LLM

請參閱[大型語言模型](#)。

較低的環境

請參閱 [環境](#)。

M

機器學習 (ML)

一種使用演算法和技術進行模式識別和學習的人工智慧。機器學習會進行分析並從記錄的資料 (例如物聯網 (IoT) 資料) 中學習，以根據模式產生統計模型。如需詳細資訊，請參閱[機器學習](#)。

主要分支

請參閱[分支](#)。

惡意軟體

旨在危及電腦安全或隱私權的軟體。惡意軟體可能會中斷電腦系統、洩露敏感資訊，或取得未經授權的存取。惡意軟體的範例包括病毒、蠕蟲、勒索軟體、特洛伊木馬程式、間諜軟體和鍵盤記錄器。

受管服務

AWS 服務 會 AWS 操作基礎設施層、作業系統和平台，而您會存取端點來存放和擷取資料。Amazon Simple Storage Service (Amazon S3) 和 Amazon DynamoDB 是受管服務的範例。這些也稱為抽象服務。

製造執行系統 (MES)

一種軟體系統，用於追蹤、監控、記錄和控制生產程序，將原物料轉換為現場成品。

MAP

請參閱[遷移加速計劃](#)。

機制

建立工具、推動工具採用，然後檢查結果以進行調整的完整程序。機制是在操作時強化和改善自身的循環。如需詳細資訊，請參閱 AWS Well-Architected Framework 中的[建置機制](#)。

成員帳戶

除了屬於組織一部分的管理帳戶 AWS 帳戶 之外的所有 AWS Organizations。一個帳戶一次只能是一個組織的成員。

製造執行系統

請參閱[製造執行系統](#)。

訊息佇列遙測傳輸 (MQTT)

根據[發佈/訂閱](#)模式的輕量型machine-to-machine(M2M) 通訊協定，適用於資源受限的 [IoT](#) 裝置。

微服務

一種小型的獨立服務，它可透過定義明確的 API 進行通訊，通常由小型獨立團隊擁有。例如，保險系統可能包含對應至業務能力 (例如銷售或行銷) 或子領域 (例如購買、索賠或分析) 的微服務。微服務的優點包括靈活性、彈性擴展、輕鬆部署、可重複使用的程式碼和適應力。如需詳細資訊，請參閱[使用無 AWS 伺服器服務整合微服務](#)。

微服務架構

一種使用獨立元件來建置應用程式的方法，這些元件會以微服務形式執行每個應用程式程序。這些微服務會使用輕量型 API，透過明確定義的介面進行通訊。此架構中的每個微服務都可以進行更新、部署和擴展，以滿足應用程式特定功能的需求。如需詳細資訊，請參閱[在上實作微服務 AWS](#)。

Migration Acceleration Program (MAP)

一種 AWS 計畫，提供諮詢支援、訓練和服務，協助組織建立強大的營運基礎，以移至雲端，並協助抵銷遷移的初始成本。MAP 包括用於有條不紊地執行舊式遷移的遷移方法以及一組用於自動化和加速常見遷移案例的工具。

大規模遷移

將大部分應用程式組合依波次移至雲端的程序，在每個波次中，都會以更快的速度移動更多應用程式。此階段使用從早期階段學到的最佳實務和經驗教訓來實作團隊、工具和流程的遷移工廠，以透過自動化和敏捷交付簡化工作負載的遷移。這是[AWS 遷移策略](#)的第三階段。

遷移工廠

可透過自動化、敏捷的方法簡化工作負載遷移的跨職能團隊。遷移工廠團隊通常包括營運、業務分析師和擁有者、遷移工程師、開發人員以及從事 Sprint 工作的 DevOps 專業人員。20% 至 50% 之間的企業應用程式組合包含可透過工廠方法優化的重複模式。如需詳細資訊，請參閱此內容集中的[遷移工廠的討論](#)和[雲端遷移工廠指南](#)。

遷移中繼資料

有關完成遷移所需的應用程式和伺服器的資訊。每種遷移模式都需要一組不同的遷移中繼資料。遷移中繼資料的範例包括目標子網路、安全群組和 AWS 帳戶。

遷移模式

可重複的遷移任務，詳細描述遷移策略、遷移目的地以及所使用的遷移應用程式或服務。範例：使用 AWS Application Migration Service 重新託管遷移至 Amazon EC2。

遷移組合評定 (MPA)

線上工具，提供驗證商業案例以遷移至的資訊 AWS 雲端。MPA 提供詳細的組合評定 (伺服器適當規模、定價、總體擁有成本比較、遷移成本分析) 以及遷移規劃 (應用程式資料分析和資料收集、應用程式分組、遷移優先順序，以及波次規劃)。[MPA 工具](#) (需要登入) 可供所有 AWS 顧問和 APN 合作夥伴顧問免費使用。

遷移準備程度評定 (MRA)

使用 AWS CAF 取得組織雲端整備狀態的洞見、識別優缺點，以及建立行動計劃以消除已識別差距的程序。如需詳細資訊，請參閱[遷移準備程度指南](#)。MRA 是 [AWS 遷移策略](#) 的第一階段。

遷移策略

用來將工作負載遷移至的方法 AWS 雲端。如需詳細資訊，請參閱本詞彙表中的 [7 個 Rs](#) 項目，並請參閱[動員您的組織以加速大規模遷移](#)。

機器學習 (ML)

請參閱[機器學習](#)。

現代化

將過時的 (舊版或單一) 應用程式及其基礎架構轉換為雲端中靈活、富有彈性且高度可用的系統，以降低成本、提高效率並充分利用創新。如需詳細資訊，請參閱 [《》中的現代化應用程式的策略 AWS 雲端](#)。

現代化準備程度評定

這項評估可協助判斷組織應用程式的現代化準備程度；識別優點、風險和相依性；並確定組織能夠在多大程度上支援這些應用程式的未來狀態。評定的結果就是目標架構的藍圖、詳細說明現代化程序的開發階段和里程碑的路線圖、以及解決已發現的差距之行動計畫。如需詳細資訊，請參閱 [《》中的評估應用程式的現代化準備 AWS 雲端](#) 程度。

單一應用程式 (單一)

透過緊密結合的程序作為單一服務執行的應用程式。單一應用程式有幾個缺點。如果一個應用程式功能遇到需求激增，則必須擴展整個架構。當程式碼庫增長時，新增或改進單一應用程式的功能也會變得更加複雜。若要解決這些問題，可以使用微服務架構。如需詳細資訊，請參閱[將單一體系分解為微服務](#)。

MPA

請參閱[遷移產品組合評估](#)。

MQTT

請參閱[訊息佇列遙測傳輸](#)。

多類別分類

一個有助於產生多類別預測的過程 (預測兩個以上的結果之一)。例如，機器學習模型可能會詢問「此產品是書籍、汽車還是電話？」或者「這個客戶對哪種產品類別最感興趣？」

可變基礎設施

更新和修改生產工作負載現有基礎設施的模型。為了提高一致性、可靠性和可預測性，AWS Well-Architected Framework 建議使用[不可變的基礎設施](#)作為最佳實務。

O

OAC

請參閱[原始存取控制](#)。

OAI

請參閱[原始存取身分](#)。

OCM

請參閱[組織變更管理](#)。

離線遷移

一種遷移方法，可在遷移過程中刪除來源工作負載。此方法涉及延長停機時間，通常用於小型非關鍵工作負載。

OI

請參閱[操作整合](#)。

OLA

請參閱[操作層級協議](#)。

線上遷移

一種遷移方法，無需離線即可將來源工作負載複製到目標系統。連接至工作負載的應用程式可在遷移期間繼續運作。此方法涉及零至最短停機時間，通常用於關鍵的生產工作負載。

OPC-UA

請參閱[開啟程序通訊 - 統一架構](#)。

開放程序通訊 - 統一架構 (OPC-UA)

用於工業自動化machine-to-machine(M2M) 通訊協定。OPC-UA 提供資料加密、身分驗證和授權機制的互通性標準。

操作水準協議 (OLA)

一份協議，闡明 IT 職能群組承諾向彼此提供的內容，以支援服務水準協議 (SLA)。

操作整備審查 (ORR)

問題及相關最佳實務的檢查清單，可協助您了解、評估、預防或減少事件和可能失敗的範圍。如需詳細資訊，請參閱 AWS Well-Architected Framework 中的[操作準備度審查 \(ORR\)](#)。

操作技術 (OT)

使用實體環境控制工業操作、設備和基礎設施的硬體和軟體系統。在製造業中，整合 OT 和資訊技術 (IT) 系統是[工業 4.0](#) 轉型的關鍵重點。

操作整合 (OI)

在雲端中將操作現代化的程序，其中包括準備程度規劃、自動化和整合。如需詳細資訊，請參閱[操作整合指南](#)。

組織追蹤

由建立的線索 AWS CloudTrail 會記錄 AWS 帳戶 組織中所有 的所有事件 AWS Organizations。在屬於組織的每個 AWS 帳戶 中建立此追蹤，它會跟蹤每個帳戶中的活動。如需詳細資訊，請參閱 CloudTrail 文件中的[建立組織追蹤](#)。

組織變更管理 (OCM)

用於從人員、文化和領導力層面管理重大、顛覆性業務轉型的架構。OCM 透過加速變更採用、解決過渡問題，以及推動文化和組織變更，協助組織為新系統和策略做好準備，並轉移至新系統和策略。在 AWS 遷移策略中，此架構稱為人員加速，因為雲端採用專案所需的變更速度。如需詳細資訊，請參閱[OCM 指南](#)。

原始存取控制 (OAC)

CloudFront 中的增強型選項，用於限制存取以保護 Amazon Simple Storage Service (Amazon S3) 內容。OAC 支援所有 S3 儲存貯體中的所有伺服器端加密 AWS KMS (SSE-KMS) AWS 區域，以及對 S3 儲存貯體的動態PUT和DELETE請求。

原始存取身分 (OAI)

CloudFront 中的一個選項，用於限制存取以保護 Amazon S3 內容。當您使用 OAI 時，CloudFront 會建立一個可供 Amazon S3 進行驗證的主體。經驗證的主體只能透過特定 CloudFront 分發來存取 S3 儲存貯體中的內容。另請參閱[OAC](#)，它可提供更精細且增強的存取控制。

ORR

請參閱[操作整備審核](#)。

OT

請參閱[操作技術](#)。

傳出 (輸出) VPC

在 AWS 多帳戶架構中，處理從應用程式內啟動之網路連線的 VPC。[AWS 安全參考架構](#)建議您使用傳入、傳出和檢查 VPC 來設定網路帳戶，以保護應用程式與更廣泛的網際網路之間的雙向介面。

P

許可界限

附接至 IAM 主體的 IAM 管理政策，可設定使用者或角色擁有的最大許可。如需詳細資訊，請參閱 IAM 文件中的[許可界限](#)。

個人身分識別資訊 (PII)

直接檢視或與其他相關資料配對時，可用來合理推斷個人身分的資訊。PII 的範例包括名稱、地址和聯絡資訊。

PII

請參閱[個人身分識別資訊](#)。

手冊

一組預先定義的步驟，可擷取與遷移關聯的工作，例如在雲端中提供核心操作功能。手冊可以採用指令碼、自動化執行手冊或操作現代化環境所需的程序或步驟摘要的形式。

PLC

請參閱[可程式設計邏輯控制器](#)。

PLM

請參閱[產品生命週期管理](#)。

政策

可定義許可的物件（請參閱[身分型政策](#)）、指定存取條件（請參閱[資源型政策](#)），或定義組織中所有帳戶的最大許可 AWS Organizations（請參閱[服務控制政策](#)）。

混合持久性

根據資料存取模式和其他需求，獨立選擇微服務的資料儲存技術。如果您的微服務具有相同的資料儲存技術，則其可能會遇到實作挑戰或效能不佳。如果微服務使用最適合其需求的資料儲存，則可以更輕鬆地實作並達到更好的效能和可擴展性。如需詳細資訊，請參閱[在微服務中啟用資料持久性](#)。

組合評定

探索、分析應用程式組合並排定其優先順序以規劃遷移的程序。如需詳細資訊，請參閱[評估遷移準備程度](#)。

述詞

傳回 true 或 的查詢條件 false，通常位於 WHERE 子句中。

述詞下推

一種資料庫查詢最佳化技術，可在傳輸前篩選查詢中的資料。這可減少必須從關聯式資料庫擷取和處理的資料量，並改善查詢效能。

預防性控制

旨在防止事件發生的安全控制。這些控制是第一道防線，可協助防止對網路的未經授權存取或不必要變更。如需詳細資訊，請參閱在 AWS 上實作安全控制中的[預防性控制](#)。

委託人

中可執行動作和存取資源 AWS 的實體。此實體通常是 AWS 帳戶、IAM 角色或使用者的根使用者。如需詳細資訊，請參閱 IAM 文件中[角色術語和概念](#)中的主體。

設計隱私權

透過整個開發程序將隱私權納入考量的系統工程方法。

私有託管區域

一種容器，它包含有關您希望 Amazon Route 53 如何回應一個或多個 VPC 內的域及其子域之 DNS 查詢的資訊。如需詳細資訊，請參閱 Route 53 文件中的[使用私有託管區域](#)。

主動控制

旨在防止部署不合規資源的[安全控制](#)。這些控制項會在佈建資源之前對其進行掃描。如果資源不符合控制項，則不會佈建。如需詳細資訊，請參閱 AWS Control Tower 文件中的[控制項參考指南](#)，並參閱實作安全[控制項中的主動](#)控制項。 AWS

產品生命週期管理 (PLM)

管理產品整個生命週期的資料和程序，從設計、開發和啟動，到成長和成熟，再到拒絕和移除。

生產環境

請參閱 [環境](#)。

可程式邏輯控制器 (PLC)

在製造中，高度可靠、可調整的電腦，可監控機器並自動化製造程序。

提示鏈結

使用一個 [LLM](#) 提示的輸出做為下一個提示的輸入，以產生更好的回應。此技術用於將複雜任務分解為子任務，或反覆精簡或展開初步回應。它有助於提高模型回應的準確性和相關性，並允許更精細、個人化的結果。

擬匿名化

將資料集中的個人識別符取代為預留位置值的程序。假名化有助於保護個人隱私權。假名化資料仍被視為個人資料。

發佈/訂閱 (pub/sub)

一種模式，可啟用微服務之間的非同步通訊，以提高可擴展性和回應能力。例如，在微服務型 [MES](#) 中，微服務可以將事件訊息發佈到其他微服務可訂閱的頻道。系統可以新增新的微服務，而無需變更發佈服務。

Q

查詢計劃

一系列步驟，如指示，用於存取 SQL 關聯式資料庫系統中的資料。

查詢計劃迴歸

在資料庫服務優化工具選擇的計畫比對資料庫環境進行指定的變更之前的計畫不太理想時。這可能因為對統計資料、限制條件、環境設定、查詢參數繫結的變更以及資料庫引擎的更新所導致。

R

RACI 矩陣

請參閱[負責、負責、諮詢、告知 \(RACI\)](#)。

RAG

請參閱[擷取增強生成](#)。

勒索軟體

一種惡意軟體，旨在阻止對計算機系統或資料的存取，直到付款為止。

RASCI 矩陣

請參閱[負責、負責、諮詢、告知 \(RACI\)](#)。

RCAC

請參閱[資料列和資料欄存取控制](#)。

僅供讀取複本

用於唯讀用途的資料庫複本。您可以將查詢路由至僅供讀取複本以減少主資料庫的負載。

重新架構師

請參閱[7 Rs](#)。

復原點目標 (RPO)

自上次資料復原點以來可接受的時間上限。這會決定最後一個復原點與服務中斷之間可接受的資料遺失。

復原時間目標 (RTO)

服務中斷和服務還原之間的可接受延遲上限。

重構

請參閱[7 個 R](#)。

區域

地理區域中的 AWS 資源集合。每個 AWS 區域 都獨立於其他，以提供容錯能力、穩定性和彈性。如需詳細資訊，請參閱[指定 AWS 區域 您的帳戶可以使用哪些](#)。

迴歸

預測數值的 ML 技術。例如，為了解決「這房子會賣什麼價格？」的問題 ML 模型可以使用線性迴歸模型，根據已知的房屋事實 (例如，平方英尺) 來預測房屋的銷售價格。

重新託管

請參閱 [7 個 R](#)。

版本

在部署程序中，它是將變更提升至生產環境的動作。

重新放置

請參閱 [7 Rs](#)。

Replatform

請參閱 [7 Rs](#)。

回購

請參閱 [7 Rs](#)。

彈性

應用程式抵禦中斷或從中斷中復原的能力。[在中規劃彈性時，高可用性和災難復原](#)是常見的考量 AWS 雲端。如需詳細資訊，請參閱[AWS 雲端 彈性](#)。

資源型政策

附接至資源的政策，例如 Amazon S3 儲存貯體、端點或加密金鑰。這種類型的政策會指定允許存取哪些主體、支援的動作以及必須滿足的任何其他條件。

負責者、當責者、事先諮詢者和事後告知者 (RACI) 矩陣

定義所有涉及遷移活動和雲端操作之各方的角色和責任的矩陣。矩陣名稱衍生自矩陣中定義的責任類型：負責人 (R)、責任 (A)、已諮詢 (C) 和知情 (I)。支援 (S) 類型為選用。如果您包含支援，則矩陣稱為 RASCI 矩陣，如果您排除它，則稱為 RACI 矩陣。

回應性控制

一種安全控制，旨在驅動不良事件或偏離安全基準的補救措施。如需詳細資訊，請參閱在 AWS 上實作安全控制中的[回應性控制](#)。

保留

請參閱 [7 Rs](#)。

淘汰

請參閱 [7 個 R](#)。

檢索增強生成 (RAG)

[一種生成式 AI](#) 技術，其中 [LLM](#) 會在產生回應之前參考訓練資料來源以外的授權資料來源。例如，RAG 模型可能會對組織的知識庫或自訂資料執行語意搜尋。如需詳細資訊，請參閱 [什麼是 RAG](#)。

輪換

定期更新 [秘密](#) 的程序，讓攻擊者更難存取登入資料。

資料列和資料欄存取控制 (RCAC)

使用已定義存取規則的基本彈性 SQL 表達式。RCAC 包含資料列許可和資料欄遮罩。

RPO

請參閱 [復原點目標](#)。

RTO

請參閱 [復原時間目標](#)。

執行手冊

執行特定任務所需的一組手動或自動程序。這些通常是為了簡化重複性操作或錯誤率較高的程序而建置。

S

SAML 2.0

許多身分提供者 (IdP) 使用的開放標準。此功能可啟用聯合單一登入 (SSO)，讓使用者可以登入 AWS Management Console 或呼叫 AWS API 操作，而無需為您組織中的每個人在 IAM 中建立使用者。如需有關以 SAML 2.0 為基礎的聯合詳細資訊，請參閱 IAM 文件中的 [關於以 SAML 2.0 為基礎的聯合](#)。

SCADA

請參閱 [監督控制和資料擷取](#)。

SCP

請參閱 [服務控制政策](#)。

秘密

您以加密形式存放的 AWS Secrets Manager 機密或限制資訊，例如密碼或使用登入資料。它由秘密值及其中繼資料組成。秘密值可以是二進位、單一字串或多個字串。如需詳細資訊，請參閱 [Secrets Manager 文件中的 Secrets Manager 秘密中的什麼內容？](#)。

依設計的安全性

透過整個開發程序將安全性納入考量的系統工程方法。

安全控制

一種技術或管理防護機制，它可預防、偵測或降低威脅行為者利用安全漏洞的能力。安全控制有四種主要類型：[預防性](#)、[偵測性](#)、[回應性](#)和[主動性](#)。

安全強化

減少受攻擊面以使其更能抵抗攻擊的過程。這可能包括一些動作，例如移除不再需要的資源、實作授予最低權限的安全最佳實務、或停用組態檔案中不必要的功能。

安全資訊與事件管理 (SIEM) 系統

結合安全資訊管理 (SIM) 和安全事件管理 (SEM) 系統的工具與服務。SIEM 系統會收集、監控和分析來自伺服器、網路、裝置和其他來源的資料，以偵測威脅和安全漏洞，並產生提醒。

安全回應自動化

預先定義和程式設計的動作，旨在自動回應或修復安全事件。這些自動化可做為[偵測或回應](#)式安全控制，協助您實作 AWS 安全最佳實務。自動化回應動作的範例包括修改 VPC 安全群組、修補 Amazon EC2 執行個體或輪換登入資料。

伺服器端加密

由接收資料的 AWS 服務 在其目的地加密資料。

服務控制政策 (SCP)

為 AWS Organizations 中的組織的所有帳戶提供集中控制許可的政策。SCP 會定義防護機制或設定管理員可委派給使用者或角色的動作限制。您可以使用 SCP 作為允許清單或拒絕清單，以指定允許或禁止哪些服務或動作。如需詳細資訊，請參閱 AWS Organizations 文件中的[服務控制政策](#)。

服務端點

的進入點 URL AWS 服務。您可以使用端點，透過程式設計方式連接至目標服務。如需詳細資訊，請參閱 AWS 一般參考 中的 [AWS 服務 端點](#)。

服務水準協議 (SLA)

一份協議，闡明 IT 團隊承諾向客戶提供的服務，例如服務正常執行時間和效能。

服務層級指標 (SLI)

服務效能方面的測量，例如其錯誤率、可用性或輸送量。

服務層級目標 (SLO)

代表服務運作狀態的目標指標，由[服務層級指標](#)測量。

共同責任模式

描述您與共同 AWS 承擔雲端安全與合規責任的模型。AWS 負責雲端的安全，而負責雲端的安全。如需詳細資訊，請參閱[共同責任模式](#)。

SIEM

請參閱[安全資訊和事件管理系統](#)。

單一故障點 (SPOF)

應用程式的單一關鍵元件故障，可能會中斷系統。

SLA

請參閱[服務層級協議](#)。

SLI

請參閱[服務層級指標](#)。

SLO

請參閱[服務層級目標](#)。

先拆分後播種模型

擴展和加速現代化專案的模式。定義新功能和產品版本時，核心團隊會進行拆分以建立新的產品團隊。這有助於擴展組織的能力和服務，提高開發人員生產力，並支援快速創新。如需詳細資訊，請參閱[中的階段式應用程式現代化方法 AWS 雲端](#)。

SPOF

請參閱[單一故障點](#)。

星狀結構描述

使用一個大型事實資料表來存放交易或測量資料的資料庫組織結構，並使用一或多個較小的維度資料表來存放資料屬性。此結構旨在用於[資料倉儲](#)或商業智慧用途。

Strangler Fig 模式

一種現代化單一系統的方法，它會逐步重寫和取代系統功能，直到舊式系統停止使用為止。此模式源自無花果藤，它長成一棵馴化樹並最終戰勝且取代了其宿主。該模式由 [Martin Fowler 引入](#)，作為重寫單一系統時管理風險的方式。如需有關如何套用此模式的範例，請參閱[使用容器和 Amazon API Gateway 逐步現代化舊版 Microsoft ASP.NET \(ASMX\) Web 服務](#)。

子網

您 VPC 中的 IP 地址範圍。子網必須位於單一可用區域。

監控控制和資料擷取 (SCADA)

在製造中，使用硬體和軟體來監控實體資產和生產操作的系統。

對稱加密

使用相同金鑰來加密及解密資料的加密演算法。

合成測試

以模擬使用者互動的方式測試系統，以偵測潛在問題或監控效能。您可以使用 [Amazon CloudWatch Synthetics](#) 來建立這些測試。

系統提示

一種向 [LLM](#) 提供內容、指示或指導方針以指示其行為的技術。系統提示有助於設定內容，並建立與使用者互動的規則。

T

標籤

做為中繼資料以組織 AWS 資源的鍵值對。標籤可協助您管理、識別、組織、搜尋及篩選資源。如需詳細資訊，請參閱[標記您的 AWS 資源](#)。

目標變數

您嘗試在受監督的 ML 中預測的值。這也被稱為結果變數。例如，在製造設定中，目標變數可能是產品瑕疵。

任務清單

用於透過執行手冊追蹤進度的工具。任務清單包含執行手冊的概觀以及要完成的一般任務清單。對於每個一般任務，它包括所需的預估時間量、擁有者和進度。

測試環境

請參閱 [環境](#)。

訓練

為 ML 模型提供資料以供學習。訓練資料必須包含正確答案。學習演算法會在訓練資料中尋找將輸入資料屬性映射至目標的模式 (您想要預測的答案)。它會輸出擷取這些模式的 ML 模型。可以使用 ML 模型，來預測您不知道的目標新資料。

傳輸閘道

可以用於互連 VPC 和內部部署網路的網路傳輸中樞。如需詳細資訊，請參閱 AWS Transit Gateway 文件中的 [什麼是傳輸閘道](#)。

主幹型工作流程

這是一種方法，開發人員可在功能分支中本地建置和測試功能，然後將這些變更合併到主要分支中。然後，主要分支會依序建置到開發環境、生產前環境和生產環境中。

受信任的存取權

將許可授予您指定的服務，以代表您在組織中 AWS Organizations 及其帳戶中執行任務。受信任的服務會在需要該角色時，在每個帳戶中建立服務連結角色，以便為您執行管理工作。如需詳細資訊，請參閱 文件中的 AWS Organizations [搭配使用 AWS Organizations 與其他 AWS 服務](#)。

調校

變更訓練程序的各個層面，以提高 ML 模型的準確性。例如，可以透過產生標籤集、新增標籤、然後在不同的設定下多次重複這些步驟來訓練 ML 模型，以優化模型。

雙比薩團隊

兩個比薩就能吃飽的小型 DevOps 團隊。雙披薩團隊規模可確保軟體開發中的最佳協作。

U

不確定性

這是一個概念，指的是不精確、不完整或未知的資訊，其可能會破壞預測性 ML 模型的可靠性。有兩種類型的不確定性：認知不確定性是由有限的、不完整的資料引起的，而隨機不確定性是由資料中固有的噪聲和隨機性引起的。如需詳細資訊，請參閱 [量化深度學習系統的不確定性](#) 指南。

未區分的任務

也稱為繁重工作，這是建立和操作應用程式的必要工作，但不為最終使用者提供直接價值或提供競爭優勢。未區分任務的範例包括採購、維護和容量規劃。

較高的環境

請參閱 [環境](#)。

V

清空

一種資料庫維護操作，涉及增量更新後的清理工作，以回收儲存並提升效能。

版本控制

追蹤變更的程序和工具，例如儲存庫中原始程式碼的變更。

VPC 對等互連

兩個 VPC 之間的連線，可讓您使用私有 IP 地址路由流量。如需詳細資訊，請參閱 Amazon VPC 文件中的 [什麼是 VPC 對等互連](#)。

漏洞

危害系統安全性的軟體或硬體瑕疵。

W

暖快取

包含經常存取的目前相關資料的緩衝快取。資料庫執行個體可以從緩衝快取讀取，這比從主記憶體或磁碟讀取更快。

暖資料

不常存取的資料。查詢這類資料時，通常可接受中等速度的查詢。

視窗函數

SQL 函數，對與目前記錄在某種程度上相關的資料列群組執行計算。視窗函數適用於處理任務，例如根據目前資料列的相對位置計算移動平均值或存取資料列的值。

工作負載

提供商業價值的資源和程式碼集合，例如面向客戶的應用程式或後端流程。

工作串流

遷移專案中負責一組特定任務的功能群組。每個工作串流都是獨立的，但支援專案中的其他工作串流。例如，組合工作串流負責排定應用程式、波次規劃和收集遷移中繼資料的優先順序。組合工作串流將這些資產交付至遷移工作串流，然後再遷移伺服器 and 應用程式。

WORM

請參閱[寫入一次，讀取許多](#)。

WQF

請參閱[AWS 工作負載資格架構](#)。

寫入一次，讀取許多 (WORM)

儲存模型，可一次性寫入資料，並防止刪除或修改資料。授權使用者可以視需要多次讀取資料，但無法變更資料。此資料儲存基礎設施被視為[不可變](#)。

Z

零時差入侵

利用[零時差漏洞](#)的攻擊，通常是惡意軟體。

零時差漏洞

生產系統中未緩解的瑕疵或漏洞。威脅行為者可以使用這種類型的漏洞來攻擊系統。開發人員經常因為攻擊而意識到漏洞。

零鏡頭提示

提供 [LLM](#) 執行任務的指示，但沒有可協助引導任務的範例 (快照)。LLM 必須使用其預先訓練的知識來處理任務。零鏡頭提示的有效性取決於任務的複雜性和提示的品質。另請參閱[少量擷取提示](#)。

殭屍應用程式

CPU 和記憶體平均使用率低於 5% 的應用程式。在遷移專案中，通常會淘汰這些應用程式。

本文為英文版的機器翻譯版本，如內容有任何歧義或不一致之處，概以英文版為準。