



Actividades de mantenimiento de las bases de datos PostgreSQL en Amazon RDS y Amazon Aurora para evitar problemas de rendimiento

AWS Guía prescriptiva



AWS Guía prescriptiva: Actividades de mantenimiento de las bases de datos PostgreSQL en Amazon RDS y Amazon Aurora para evitar problemas de rendimiento

Table of Contents

Introducción	1
Resultados empresariales específicos	1
Control de simultaneidad multiversión (MVCC)	1
Aspirar y analizar mesas automáticamente	3
Parámetros relacionados con la memoria del aspirador automático	5
Ajuste de los parámetros del aspirador automático	6
A nivel de clúster o instancia	6
Nivel de tabla	6
Utilice ajustes de aspiración automática agresivos a nivel de mesa	6
Ventajas y limitaciones	8
Aspirar y analizar mesas manualmente	9
Ejecutar las operaciones de aspiración y limpieza en paralelo	10
Reescribir una tabla completa con VACUUM FULL	14
Eliminar la hinchazón con pg_repack	15
Reconstrucción de índices	17
Crear un índice nuevo	18
Reconstruir un índice	18
Ejemplos	19
Ejemplo: Recuperación de espacio mediante autovacuum y VACUUM FULL	21
Recursos	26
Historial de documentos	27
Glosario	28
#	28
A	29
B	32
C	34
D	37
E	42
F	44
G	46
H	47
I	49
L	51
M	52

O	57
P	60
Q	63
R	63
S	66
T	70
U	72
V	72
W	73
Z	74
.....	lxxv

Actividades de mantenimiento de las bases de datos PostgreSQL en Amazon RDS y Amazon Aurora para evitar problemas de rendimiento

Anuradha Chintla, Rajesh Madiwale y Srinivas Potlacheruvu, Amazon Web Services (AWS)

Diciembre de 2023 ([historial de documentos](#))

Amazon Aurora PostgreSQL Compatible Edition y Amazon Relational Database Service (Amazon RDS) para PostgreSQL son servicios de bases de datos relacionales totalmente gestionados para bases de datos PostgreSQL. Estos servicios gestionados liberan al administrador de la base de datos de muchas tareas de mantenimiento y administración. Sin embargo, algunas tareas de mantenimiento, por ejemplo VACUUM, requieren una estrecha supervisión y configuración en función del uso de la base de datos. Esta guía describe las actividades de mantenimiento de PostgreSQL en Amazon RDS y Aurora.

Resultados empresariales específicos

El rendimiento de las bases de datos es una medida clave que sustenta el éxito de una empresa. La realización de actividades de mantenimiento en sus bases de datos compatibles con Aurora PostgreSQL y de Amazon RDS for PostgreSQL ofrece las siguientes ventajas:

- Ayuda a lograr un rendimiento óptimo de las consultas
- Libera espacio abultado para que lo reutilicen en futuras transacciones
- Evita el envolvimiento de las transacciones
- Ayuda al optimizador a generar buenos planes
- Garantiza un uso adecuado del índice

Control de simultaneidad multiversión (MVCC)

El mantenimiento de la base de datos PostgreSQL requiere comprender el control de concurrencia multiversión (MVCC), que es un mecanismo de PostgreSQL. Cuando se procesan varias transacciones simultáneamente en la base de datos, el MVCC garantiza que se mantengan la atomicidad y el aislamiento, que son dos características de las transacciones de atomicidad,

consistencia, aislamiento y durabilidad (ACID). En el MVCC, cada operación de escritura genera una nueva versión de los datos y almacena la versión anterior. Los lectores y los escritores no se bloquean entre sí. Cuando una transacción lee datos, el sistema elige una de las versiones para aislar la transacción. PostgreSQL y algunas bases de datos relacionales utilizan una adaptación de MVCC denominada aislamiento de instantáneas (SI). Por ejemplo, Oracle implementa la SI mediante segmentos de reversión. Durante una operación de escritura, Oracle escribe la versión anterior de los datos en un segmento de reversión y sobrescribe el área de datos con la nueva versión. Las bases de datos PostgreSQL implementan la SI mediante reglas de verificación de visibilidad para evaluar las versiones. Cuando se colocan datos nuevos en una página de tabla, PostgreSQL utiliza estas reglas para seleccionar la versión adecuada de los datos para una operación de lectura.

Al modificar los datos de una fila de la tabla, PostgreSQL utiliza MVCC para mantener varias versiones de la fila. Durante UPDATE las DELETE operaciones de la tabla, la base de datos conserva las versiones antiguas de las filas para otras transacciones en ejecución que puedan necesitar una vista coherente de los datos. Estas versiones antiguas se denominan filas muertas (tuplas). Una colección de tuplas muertas produce hinchazón. Una gran cantidad de espacio en la base de datos puede provocar una serie de problemas, como una generación deficiente del plan de consultas, un rendimiento lento de las consultas y un mayor uso del espacio en disco para almacenar las versiones anteriores.

Eliminar la sobrecarga y mantener una base de datos en buen estado requiere un mantenimiento periódico, que incluye estas actividades, que se describen en las siguientes secciones:

- [Aspirar y analizar las mesas automáticamente](#)
- [Aspirar y analizar las mesas manualmente](#)
- [Eliminar la hinchazón con `pg\_repack`](#)
- [Reconstrucción de índices](#)

Aspirar y analizar mesas automáticamente

Autovacuum es un demonio (es decir, se ejecuta en segundo plano) que aspira (limpia) automáticamente las tuplas muertas, recupera el almacenamiento y recopila estadísticas. Comprueba si hay tablas hinchadas en la base de datos y las borra para reutilizar el espacio. Supervisa las tablas y los índices de las bases de datos y los añade a una tarea vacía cuando alcanzan un límite específico de operaciones de actualización o eliminación.

Autovacuum gestiona la aspiración automatizando VACUUM PostgreSQL y los comandos. ANALYZE VACUUM elimina la sobrecarga de las tablas y recupera espacio, a la vez que ANALYZE actualiza las estadísticas que permiten al optimizador elaborar planes eficientes. VACUUM también lleva a cabo una importante tarea denominada congelación por vacío para evitar problemas relacionados con los identificadores de las transacciones en la base de datos. Cada fila que se actualiza en la base de datos recibe un ID de transacción del mecanismo de control de transacciones de PostgreSQL. Estos IDs controlan la visibilidad de la fila frente a otras transacciones simultáneas. El identificador de la transacción es un número de 32 bits. Dos mil millones IDs siempre se guardan en el pasado visible. Los restantes (unos 2.200 millones) IDs se conservan para transacciones que se realizarán en el futuro y se ocultan de la transacción actual. PostgreSQL requiere una limpieza y congelación ocasionales de las filas antiguas para evitar que las transacciones se acumulen y hagan invisibles las filas antiguas existentes cuando se crean nuevas transacciones. Para obtener más información, consulte [Preventing Transaction ID Wraparound Failures](#) en la documentación de PostgreSQL.

Se recomienda utilizar el aspirador automático, que está activado de forma predeterminada. Entre sus parámetros se incluyen los siguientes.

Parámetro	Descripción	Predeterminado para Amazon RDS	Predeterminado para Aurora
autovacuum_vacuum_threshold	El número mínimo de operaciones de actualización o eliminación de tuplas que deben realizarse en una mesa antes de que la aspiradora automática la aspire.	50 operaciones	50 operaciones

autovacuum_analyze_threshold	El número mínimo de inserciones, actualizaciones o eliminaciones de tuplas que deben producirse en una tabla antes de que Autovacuum la analice.	50 operaciones	50 operaciones
autovacuum_vacuum_scale_factor	El porcentaje de tuplas que deben modificarse en una tabla antes de aspirarla con una aspiradora automática.	0,2%	0,1%
autovacuum_analyze_scale_factor	El porcentaje de tuplas que deben modificarse en una tabla antes de que la aspiradora automática la analice.	0,05%	0,05%
autovacuum_freeze_max_age	La edad máxima que debe permanecer inmovilizada IDs antes de que se vacíe una mesa para evitar problemas relacionados con el identificador de la transacción.	200 000 000 de transacciones	200 000 000 de transacciones

Autovacuum crea una lista de tablas para procesarlas en función de fórmulas de umbral específicas, de la siguiente manera.

- Umbral para correr VACUUM sobre una mesa:

```
vacuum threshold = autovacuum_vacuum_threshold + (autovacuum_vacuum_scale_factor *  
Total row count of table)
```

- Umbral para correr ANALYZE sobre una mesa:

```
analyze threshold = autovacuum_analyze_threshold + (autovacuum_analyze_scale_factor *  
Total row count of table)
```

Para tablas pequeñas o medianas, los valores predeterminados pueden ser suficientes. Sin embargo, una tabla grande que tenga modificaciones de datos frecuentes tendrá un mayor número de tuplas inactivas. En este caso, es posible que la aspiradora automática procese la tabla con frecuencia para realizar tareas de mantenimiento y que el mantenimiento de otras tablas se retrase o se ignore hasta que se acabe la tabla grande. Para evitarlo, puede ajustar los parámetros de la aspiradora automática que se describen en la siguiente sección.

Parámetros relacionados con la memoria del aspirador automático

autovacuum_max_workers

Especifica el número máximo de procesos de aspiración automática (distintos del lanzador de aspiración automática) que se pueden ejecutar al mismo tiempo. Este parámetro solo se puede configurar al iniciar el servidor. Si el proceso de aspiración automática está ocupado con una tabla grande, este parámetro ayuda a limpiar otras tablas.

maintenance_work_mem

Especifica la cantidad máxima de memoria que deben utilizar las operaciones de mantenimiento VACUUM, como CREATE INDEX, y ALTER. En Amazon RDS y Aurora, la memoria se asigna en función de la clase de instancia mediante la fórmula $\text{GREATEST}(\{\text{DBInstanceClassMemory}/63963136 \times 1024\}, 65536)$. Cuando se ejecuta la aspiradora automática, se puede asignar hasta 3 `autovacuum_max_workers` veces el valor calculado, así que procure no establecer el valor demasiado alto. Para controlar esto, puede configurarlo `autovacuum_work_mem` por separado.

autovacuum_work_mem

Especifica la cantidad máxima de memoria que utilizará cada proceso de trabajo de autovacuum. El valor predeterminado de este parámetro es -1, lo que indica que debe utilizar el valor de `maintenance_work_mem` en su lugar.

Para obtener más información sobre los parámetros de memoria del aspirador automático, consulte [Asignación de memoria para el aspirador automático en la documentación](#) de Amazon RDS.

Ajuste de los parámetros del aspirador automático

Es posible que los usuarios necesiten ajustar los parámetros del aspirador automático en función de sus operaciones de actualización y eliminación. La configuración de los siguientes parámetros se puede establecer a nivel de tabla, instancia o clúster.

A nivel de clúster o instancia

Como ejemplo, veamos una base de datos bancaria en la que se esperan operaciones de lenguaje de manipulación continua de datos (DML). Para mantener el estado de la base de datos, debe ajustar los parámetros de autovacuum a nivel de clúster para Aurora y a nivel de instancia para Amazon RDS, y aplicar también el mismo grupo de parámetros al lector. En el caso de una conmutación por error, se deberían aplicar los mismos parámetros al nuevo grabador.

Nivel de tabla

Por ejemplo, en una base de datos para la entrega de alimentos en la que se esperan operaciones de DML continuas en una sola tabla llamada `orders`, debería plantearse ajustar el `autovacuum_analyze_threshold` parámetro a nivel de tabla mediante el siguiente comando:

```
ALTER TABLE <table_name> SET (autovacuum_analyze_threshold = <threshold rows>)
```

Utilice ajustes de aspiración automática agresivos a nivel de mesa

La `orders` tabla de ejemplo que tiene operaciones de actualización y eliminación continuas se convierte en una buena opción para aspirar debido a la configuración de aspiración automática predeterminada. Esto lleva a que se generen planes incorrectos y a que las consultas sean lentas. Eliminar la sobrecarga y actualizar las estadísticas requiere una configuración de autovació agresiva a nivel de tabla.

Para determinar la configuración, lleve un registro de la duración de las consultas que se ejecutan en esta tabla e identifique el porcentaje de operaciones de DML que dan lugar a cambios en el plan.

La `pg_stat_all_table` vista le ayuda a realizar un seguimiento de las operaciones de inserción, actualización y eliminación.

Supongamos que el optimizador genera planes incorrectos cada vez que cambia el 5 por ciento de la `orders` tabla. En este caso, debe cambiar el umbral al 5 por ciento de la siguiente manera:

```
ALTER TABLE orders SET (autovacuum_analyze_threshold = 0.05 and  
autovacuum_vacuum_threshold = 0.05)
```

 Tip

Seleccione cuidadosamente los ajustes de aspiración automática agresivos para evitar un alto consumo de recursos.

Para obtener más información, consulte los siguientes temas:

- [Descripción del autovacuum en entornos Amazon RDS para AWS PostgreSQL](#) (entrada del blog)
- [Aspiración automática \(documentación de PostgreSQL\)](#)
- [Ajuste de los parámetros de PostgreSQL en Amazon RDS y Amazon AWS Aurora](#) (guía prescriptiva)

Para asegurarse de que la aspiradora automática funciona de forma eficaz, supervise las filas muertas, el uso del disco y la última vez que la aspiradora automática o ANALYZE se ejecutó de forma regular. La `pg_stat_all_tables` vista proporciona información sobre cada tabla (`relname`) y sobre cuántas tuplas muertas (`n_dead_tup`) hay en la tabla.

Supervisar el número de tuplas inactivas en cada tabla, especialmente en las que se actualizan con frecuencia, permite determinar si los procesos de aspiración automática eliminan periódicamente las tuplas inactivas para reutilizar su espacio en disco y mejorar el rendimiento. Puede utilizar la siguiente consulta para comprobar el número de tuplas inactivas y cuándo se utilizó la última aspiradora automática en las tablas:

```
SELECT  
relname AS TableName, n_live_tup AS LiveTuples, n_dead_tup AS DeadTuples,  
last_autovacuum AS Autovacuum, last_autoanalyze AS Autoanalyze  
FROM  
pg_all_user_tables;
```

Ventajas y limitaciones

La aspiradora automática ofrece las siguientes ventajas:

- Elimina automáticamente la hinchazón de las mesas.
- Evita que la ID de la transacción se distorsione.
- Mantiene actualizadas las estadísticas de la base de datos.

Limitaciones:

- Si las consultas utilizan el procesamiento paralelo, es posible que la cantidad de procesos de trabajo no sea suficiente para el autovacuum.
- Si la aspiradora automática se ejecuta durante las horas pico, es posible que aumente la utilización de los recursos. Debe ajustar los parámetros para solucionar este problema.
- Si las páginas de la tabla están ocupadas en otra sesión, es posible que Autovacuum las omita.
- Autovacuum no puede acceder a las tablas temporales.

Aspirar y analizar mesas manualmente

Si su base de datos se vacía mediante el proceso de vacío automático, se recomienda evitar ejecutar aspiraciones manuales en toda la base de datos con demasiada frecuencia. Un vacío manual puede provocar cargas de E/S innecesarias o picos de CPU y, además, es posible que no se eliminen las tuplas inactivas. Utilice las aspiradoras manuales table-by-table únicamente si es realmente necesario, por ejemplo, cuando la proporción de tuplas activas y muertas sea baja o cuando haya espacios largos entre las aspiradoras automáticas. Además, debe utilizar aspiradoras manuales cuando la actividad del usuario sea mínima.

Autovacuum también mantiene actualizadas las estadísticas de una tabla. Al ejecutar el ANALYZE comando manualmente, reconstruye estas estadísticas en lugar de actualizarlas. La reconstrucción de las estadísticas cuando ya están actualizadas mediante el proceso normal de aspiración automática puede provocar la utilización de los recursos del sistema.

Se recomienda ejecutar los comandos [VACUUM](#) y [ANALYZE](#) manualmente en los siguientes escenarios:

- Durante las horas de menor actividad y en las mesas más concurridas, es posible que la aspiración automática no sea suficiente.
- Inmediatamente después de cargar datos de forma masiva en la tabla de destino. En este caso, si se ejecuta ANALYZE manualmente, se regeneran completamente las estadísticas, lo cual es una mejor opción que esperar a que comience la aspiradora automática.
- Para aspirar tablas temporales (la aspiradora automática no puede acceder a ellas).

Para reducir el impacto de E/S al ejecutar los ANALYZE comandos VACUUM y en la actividad simultánea de la base de datos, puede utilizar el parámetro. `vacuum_cost_delay` En muchas situaciones, los comandos de mantenimiento, por ejemplo, ANALYZE no tienen por qué finalizarse rápidamente. VACUUM Sin embargo, estos comandos no deberían interferir con la capacidad del sistema para realizar otras operaciones de la base de datos. Para evitarlo, puede activar los retardos de vacío basados en los costes mediante el uso del `vacuum_cost_delay` parámetro. Este parámetro está desactivado de forma predeterminada para los VACUUM comandos ejecutados manualmente. Para activarlo, configúrelo en un valor distinto de cero.

Ejecutar las operaciones de aspiración y limpieza en paralelo

La opción [PARALLEL](#) del VACUUM comando utiliza trabajadores paralelos para las fases de aspiración y limpieza de índices y está deshabilitada de forma predeterminada. El número de trabajadores paralelos (el grado de paralelismo) viene determinado por el número de índices de la tabla y el usuario puede especificarlo. Si ejecuta VACUUM operaciones paralelas sin un argumento entero, el grado de paralelismo se calcula en función del número de índices de la tabla.

Los siguientes parámetros le ayudan a configurar la aspiración paralela en Amazon RDS para que sean compatibles con PostgreSQL y Aurora PostgreSQL:

- `max_worker_processes` establece el [número máximo](#) de procesos de trabajo simultáneos.
- [min_parallel_index_scan_size](#) establece la cantidad mínima de datos de índice que se deben escanear para poder considerar un escaneo paralelo.
- [max_parallel_maintenance_workers](#) establece el número máximo de trabajadores paralelos que se pueden iniciar con un único comando de utilidad.

Note

La opción solo se utiliza para aspirar `PARALLEL`. No afecta al comando `ANALYZE`.

El siguiente ejemplo ilustra el comportamiento de la base de datos cuando se usa de forma manual VACUUM y ANALYZE en una base de datos.

Esta es una tabla de ejemplo en la que se ha desactivado la aspiradora automática (solo con fines ilustrativos; no se recomienda deshabilitar la aspiradora automática):

```
create table t1 ( a int, b int, c int );
alter table t1 set (autovacuum_enabled=false);
```

```
apgl=> \d+ t1
Table "public.t1"
Column | Type | Collation | Nullable | Default | Storage | Stats target | Description
-----+-----+-----+-----+-----+-----+-----+-----
a | integer | | | plain | | | 
b | integer | | | plain | | |
```

```
c | integer | | | plain | |
Access method: heap
Options: autovacuum_enabled=false
```

Añada 1 millón de filas a la tabla t1:

```
apgl=> select count(*) from t1;
count
1000000
(1 row)
```

Estadísticas de la tabla t1:

```
select * from pg_stat_all_tables where relname='t1';
-[ RECORD 1 ]-----+-----
relid          | 914744
schemaname     | public
relname        | t1
seq_scan       | 0
seq_tup_read   | 0
idx_scan       | 
idx_tup_fetch  | 
n_tup_ins      | 1000000
n_tup_upd      | 0
n_tup_del      | 0
n_tup_hot_upd  | 0
n_live_tup     | 1000000
n_dead_tup     | 0
n_mod_since_analyze | 1000000
last_vacuum    | 
last_autovacuum | 
last_analyze   | 
last_autoanalyze | 
vacuum_count   | 0
autovacuum_count | 0
analyze_count  | 0
autoanalyze_count | 0
```

Añada un índice:

```
create index i2 on t1 (b,a);
```

Ejecute el EXPLAIN comando (Plan 1):

```
Bitmap Heap Scan on t1 (cost=10521.17..14072.67 rows=5000 width=4)
Recheck Cond: (a = 5)
# Bitmap Index Scan on i2 (cost=0.00..10519.92 rows=5000 width=0)
Index Cond: (a = 5)
(4 rows)
```

Ejecute el EXPLAIN ANALYZE comando (Plan 2):

```
explain (analyze, buffers, costs off) select a from t1 where b = 5;
QUERY PLAN
Bitmap Heap Scan on t1 (actual time=0.023..0.024 rows=1 loops=1)
Recheck Cond: (b = 5)
Heap Blocks: exact=1
Buffers: shared hit=4
# Bitmap Index Scan on i2 (actual time=0.016..0.016 rows=1 loops=1)
Index Cond: (b = 5)
Buffers: shared hit=3
Planning Time: 0.054 ms
Execution Time: 0.076 ms
(9 rows)
```

Los EXPLAIN ANALYZE comandos EXPLAIN y muestran planes diferentes, ya que la aspiradora automática estaba desactivada en la tabla y el ANALYZE comando no se ejecutaba manualmente. Ahora vamos a actualizar un valor de la tabla y regenerar el EXPLAIN ANALYZE plan:

```
update t1 set a=8 where b=5;
explain (analyze, buffers, costs off) select a from t1 where b = 5;
```

El EXPLAIN ANALYZE comando (Plan 3) ahora muestra:

```
apgl=> explain (analyze, buffers, costs off) select a from t1 where b = 5;
QUERY PLAN
Bitmap Heap Scan on t1 (actual time=0.075..0.076 rows=1 loops=1)
Recheck Cond: (b = 5)
Heap Blocks: exact=1
Buffers: shared hit=5
# Bitmap Index Scan on i2 (actual time=0.017..0.017 rows=2 loops=1)
Index Cond: (b = 5)
Buffers: shared hit=3
```



```
Planning Time: 0.053 ms
Execution Time: 0.125 ms
```

Si comparas los costos entre el Plan 2 y el Plan 3, verás las diferencias en el tiempo de planificación y ejecución, ya que aún no hemos recopilado estadísticas.

Ahora pongamos un manual ANALYZE sobre la mesa y, a continuación, comprobemos las estadísticas y regeneremos el plan:

```
apgl=> analyze t1
apgl# ;
ANALYZE
Time: 212.223 ms

apgl=> select * from pg_stat_all_tables where relname='t1';
-[ RECORD 1 ]-----+-----
relid          | 914744
schemaname     | public
relname        | t1
seq_scan       | 3
seq_tup_read   | 1000000
idx_scan       | 3
idx_tup_fetch  | 3
n_tup_ins      | 1000000
n_tup_upd      | 1
n_tup_del      | 0
n_tup_hot_upd  | 0
n_live_tup     | 1000000
n_dead_tup     | 1
n_mod_since_analyze | 0
last_vacuum    |
last_autovacuum |
last_analyze   | 2023-04-15 11:39:02.075089+00
last_autoanalyze |
vacuum_count   | 0
autovacuum_count | 0
analyze_count  | 1
autoanalyze_count | 0

Time: 148.347 ms
```

Ejecute el EXPLAIN ANALYZE comando (Plan 4):

```
apgl=> explain (analyze,buffers,costs off) select a from t1 where b = 5;
QUERY PLAN
Index Only Scan using i2 on t1 (actual time=0.022..0.023 rows=1 loops=1)
Index Cond: (b = 5)
Heap Fetches: 1
Buffers: shared hit=4
Planning Time: 0.056 ms
Execution Time: 0.068 ms
(6 rows)

Time: 138.462 ms
```

Si compara todos los resultados del plan después de analizar manualmente la tabla y recopilar estadísticas, observará que el Plan 4 del optimizador es mejor que los demás y, además, reduce el tiempo de ejecución de las consultas. En este ejemplo, se muestra la importancia de ejecutar actividades de mantenimiento en la base de datos.

Reescribir una tabla completa con VACUUM FULL

Al ejecutar el `VACUUM` comando con el `FULL` parámetro, se reescribe todo el contenido de una tabla en un nuevo archivo de disco sin espacio adicional y se devuelve al sistema operativo el espacio no utilizado. Esta operación es mucho más lenta y requiere `ACCESS EXCLUSIVE` bloquear cada tabla. También requiere espacio adicional en el disco, ya que escribe una nueva copia de la tabla y no libera la copia anterior hasta que se completa la operación.

`VACUUM FULL` puede resultar útil en los siguientes casos:

- Cuando desee recuperar una cantidad significativa de espacio de las mesas.
- Cuando desee recuperar espacio adicional en tablas que no son de clave principal.

Le recomendamos que las utilice `VACUUM FULL` cuando tenga tablas de claves no principales, si la base de datos puede tolerar el tiempo de inactividad.

Como `VACUUM FULL` requiere más bloqueos que otras operaciones, resulta más caro ejecutarla en bases de datos cruciales. Para reemplazar este método, puede usar la `pg_repack` extensión, que se describe en la [siguiente sección](#). Esta opción es similar a `VACUUM FULL`, pero requiere un bloqueo mínimo, y es compatible con Amazon RDS for PostgreSQL y Aurora, compatible con PostgreSQL.

Eliminar la hinchazón con pg_repack

Puede usar la `pg_repack` extensión para eliminar el exceso de tablas e índices con un bloqueo mínimo de la base de datos. Puede crear esta extensión en la instancia de base de datos y ejecutar el `pg_repack` cliente (si la versión del cliente coincide con la versión de la extensión) desde Amazon Elastic Compute Cloud (Amazon EC2) o desde un ordenador que se pueda conectar a la base de datos.

Por el contrario `VACUUM FULL`, `pg_repack` no requiere tiempo de inactividad ni un período de mantenimiento, y no bloquea otras sesiones.

`pg_repack` es útil en situaciones en las que `VACUUM FULL` o `REINDEX` puede que no, funcione. Crea una tabla nueva que contiene los datos de la tabla hinchada, realiza un seguimiento de los cambios con respecto a la tabla original y, a continuación, reemplaza la tabla original por la nueva. No bloquea la tabla original para operaciones de lectura o escritura mientras crea la nueva tabla.

Se puede utilizar `pg_repack` para una tabla completa o para un índice. Para ver una lista de tareas, consulta la documentación de [pg_repack](#).

Limitaciones:

- Para que se ejecute `pg_repack`, la tabla debe tener una clave principal o un índice único.
- `pg_repack` no funcionará con tablas temporales.
- `pg_repack` no funcionará en tablas que tengan índices globales.
- Cuando `pg_repack` está en curso, no se pueden realizar operaciones DDL en las tablas.

En la siguiente tabla se describen las diferencias entre `pg_repack` y `VACUUM FULL`.

<code>VACUUM FULL</code>	<code>pg_repack</code>
Comando integrado	Una extensión que ejecutas desde Amazon EC2 o desde tu ordenador local
Requiere un <code>ACCESS EXCLUSIVE</code> candado mientras funciona en una mesa	Requiere un <code>ACCESS EXCLUSIVE</code> candado solo por un período breve

Funciona con todas las mesas

Funciona solo en tablas que tienen claves principales y únicas

Requiere el doble del espacio de almacenamiento que consumen la tabla y los índices

Requiere el doble del almacenamiento que consumen la tabla y los índices

Para ejecutar `pg_repack` en una tabla, usa el comando:

```
pg_repack -h <host> -d <dbname> --table <tablename> -k
```

Para ejecutar `pg_repack` en un índice, utilice el comando:

```
pg_repack -h <host> -d <dbname> --index <index name>
```

Para obtener más información, consulte la entrada del AWS blog [Eliminar la hinchazón de Amazon Aurora y RDS para PostgreSQL](#) con `pg_repack`.

Reconstrucción de índices

El comando REINDEX de [PostgreSQL](#) reconstruye un índice utilizando los datos almacenados en la tabla del índice y reemplazando la copia anterior del índice. Le recomendamos que lo utilice REINDEX en los siguientes escenarios:

- Cuando un índice se daña y ya no contiene datos válidos. Esto puede ocurrir como resultado de errores de software o hardware.
- Cuando las consultas que anteriormente utilizaban el índice dejan de usarlo.
- Cuando el índice se llena con un gran número de páginas vacías o casi vacías. Se debe ejecutar REINDEX cuando el porcentaje de hinchazón (`bloat_pct`) sea superior a 20.

Las páginas de índice que están completamente vacías se recuperan para su reutilización. Sin embargo, recomendamos volver a indexarlas periódicamente si se han eliminado las claves de índice de una página pero se mantiene el espacio asignado.

Volver a crear el índice ayuda a mejorar el rendimiento de las consultas. Puede volver a crear un índice de tres maneras, tal y como se describe en la siguiente tabla.

Método	Descripción	Limitaciones
CREATE INDEXy DROP INDEX con la opción CONCURRENTLY	Crea un índice nuevo y elimina el índice anterior. El optimizador genera planes utilizando el índice recién creado en lugar del índice anterior. Durante las horas de menor actividad, puede eliminar el índice anterior.	La creación del índice lleva más tiempo cuando se utiliza CONCURRENTLY esta opción, ya que tiene que realizar un seguimiento de todos los cambios entrantes. Cuando los cambios están bloqueados, el proceso se marca como completado.
REINDEXcon la CONCURRENTLY opción	Bloquea las operaciones de escritura durante el proceso de reconstrucción. La versión 12 y las versiones posteriores de PostgreSQL ofrecen esta	El uso CONCURRENTLY requiere más tiempo para reconstruir el índice.

opción, lo que CONCURRENTLY evita estos bloqueos.

pg_repack extensión

Limpia la hinchazón de una tabla y reconstruye el índice.

Debe ejecutar esta extensión desde una EC2 instancia o desde un equipo local que esté conectado a la base de datos.

Crear un índice nuevo

Los CREATE INDEX comandos DROP INDEX y, cuando se utilizan juntos, reconstruyen un índice:

```
DROP INDEX <index_name>
CREATE INDEX <index_name> ON TABLE <table_name> (<column1>[,<column2>])
```

La desventaja de este enfoque es su exclusividad, lo que afecta al rendimiento durante esta actividad. El DROP INDEX comando adquiere un bloqueo exclusivo, que bloquea las operaciones de lectura y escritura en la tabla. El CREATE INDEX comando bloquea las operaciones de escritura en la tabla. Permite operaciones de lectura, pero son caras durante la creación del índice.

Reconstruir un índice

El REINDEX comando le ayuda a mantener un rendimiento uniforme de la base de datos. Al realizar un gran número de operaciones de DML en una tabla, se produce una saturación tanto de la tabla como de los índices. Los índices se utilizan para acelerar las búsquedas en las tablas y mejorar el rendimiento de las consultas. La saturación de los índices afecta al rendimiento de las búsquedas y las consultas. Por lo tanto, se recomienda volver a indexar las tablas que tienen un gran volumen de operaciones de DML para mantener la coherencia en el rendimiento de las consultas.

El REINDEX comando reconstruye el índice desde cero al bloquear las operaciones de escritura en la tabla subyacente, pero permite las operaciones de lectura en la tabla. Sin embargo, bloquea las operaciones de lectura en el índice. Las consultas que utilizan el índice correspondiente están bloqueadas, pero las demás consultas no.

La versión 12 de PostgreSQL introdujo un nuevo parámetro opcional CONCURRENTLY, que reconstruye el índice desde cero, pero no bloquea las operaciones de escritura o lectura en la tabla

ni en las consultas que utilizan el índice. Sin embargo, se tarda más tiempo en completar el proceso cuando se utiliza esta opción.

Ejemplos

Crear y eliminar un índice

Cree un índice nuevo con la CONCURRENTLY opción:

```
create index CONCURRENTLY on table(columns) ;
```

Elimine el índice anterior con la CONCURRENTLY opción:

```
drop index CONCURRENTLY <index name> ;
```

Reconstruir un índice

Para volver a crear un índice único:

```
reindex index <index name> ;
```

Para volver a generar todos los índices de una tabla:

```
reindex table <table name> ;
```

Para reconstruir todos los índices de un esquema:

```
reindex schema <schema name> ;
```

Reconstruir un índice de forma simultánea

Para volver a crear un índice único:

```
reindex index CONCURRENTLY <indexname> ;
```

Para volver a generar todos los índices de una tabla:

```
reindex table CONCURRENTLY <tablename> ;
```

Para reconstruir todos los índices de un esquema:

```
reindex schema CONCURRENTLY <schemaname> ;
```

Reconstruir o reubicar índices únicamente

Para volver a crear un índice único:

```
pg_repack -h <hostname> -d <dbname> -i <indexname> -k
```

Para volver a crear todos los índices:

```
pg_repack -h <hostname> -d <dbname> -x <indexname> -t <tablename> -k
```


Ejemplo: Recuperación de espacio mediante autovacuum y VACUUM FULL

Como ejemplo, creemos una emp tabla de tablas con 500 000 filas y, a continuación, actualicemos las filas con nuevos valores. El autovacuum está activado, por lo que ejecutará ambos VACUUM ANALYZE comandos de esta tabla para eliminar la hinchazón y recuperar espacio. El espacio recuperado se puede reutilizar, pero no se devolverá al sistema operativo.

La siguiente consulta determina la hinchazón de la tabla:

```
-- WARNING: When run with a non-superuser role, the query inspects only indexes on
-- tables you are granted to read.
-- WARNING: Rows with is_na = 't' are known to have bad statistics ("name" type is not
-- supported).
-- This query is compatible with PostgreSQL 8.2 and later.
SELECT current_database(), nspname AS schemaname, tblname, idxname,
bs*(relpages)::bigint AS real_size,
bs*(relpages-est_pages)::bigint AS extra_size,
100 * (relpages-est_pages)::float / relpages AS extra_pct,
fillfactor,
CASE WHEN relpages > est_pages_ff
THEN bs*(relpages-est_pages_ff)
ELSE 0
END AS bloat_size,
100 * (relpages-est_pages_ff)::float / relpages AS bloat_pct,
is_na
-- , 100-(pst).avg_leaf_density AS pst_avg_bloat, est_pages, index_tuple_hdr_bm,
maxalign, pagehdr, nulldatawidth, nulldatahdrwidth, reltuples, relpages -- (DEBUG
INFO)
FROM (
SELECT coalesce(1 +
ceil(reltuples/floor((bs-pageopqdata-pagehdr)/(4+nulldatahdrwidth)::float)), 0
-- ItemIdData size + computed avg size of a tuple (nulldatahdrwidth)
) AS est_pages,
coalesce(1 +
ceil(reltuples/floor((bs-pageopqdata-pagehdr)*fillfactor/
(100*(4+nulldatahdrwidth)::float))), 0
) AS est_pages_ff,
bs, nspname, tblname, idxname, relpages, fillfactor, is_na
```

```

-- , pgstatindex(idxoid) AS pst, index_tuple_hdr_bm, maxalign, pagehdr,
nulldatawidth, nulldatahdrwidth, reltuples -- (DEBUG INFO)
FROM (
    SELECT maxalign, bs, nspname, tblname, idxname, reltuples, relpages, idxoid,
    fillfactor,
        ( index_tuple_hdr_bm +
            maxalign - CASE -- Add padding to the index tuple header to align on
MAXALIGN
                WHEN index_tuple_hdr_bm%maxalign = 0 THEN maxalign
                ELSE index_tuple_hdr_bm%maxalign
            END
        + nulldatawidth + maxalign - CASE -- Add padding to the data to align on
MAXALIGN
                WHEN nulldatawidth = 0 THEN 0
                WHEN nulldatawidth::integer%maxalign = 0 THEN maxalign
                ELSE nulldatawidth::integer%maxalign
            END
        )::numeric AS nulldatahdrwidth, pagehdr, pageopqdata, is_na
    -- , index_tuple_hdr_bm, nulldatawidth -- (DEBUG INFO)
FROM (
    SELECT n.nspname, i.tblname, i.idxname, i.reltuples, i.relpages,
        i.idxoid, i.fillfactor, current_setting('block_size')::numeric AS bs,
        CASE -- MAXALIGN: 4 on 32bits, 8 on 64bits (and mingw32 ?)
            WHEN version() ~ 'mingw32' OR version() ~ '64-bit|x86_64|ppc64|ia64|
amd64' THEN 8
            ELSE 4
        END AS maxalign,
        /* per page header, fixed size: 20 for 7.X, 24 for others */
        24 AS pagehdr,
        /* per page btree opaque data */
        16 AS pageopqdata,
        /* per tuple header: add IndexAttributeBitMapData if some cols are null-
able */
        CASE WHEN max(coalesce(s.null_frac,0)) = 0
            THEN 8 -- IndexTupleData size
            ELSE 8 + (( 32 + 8 - 1 ) / 8) -- IndexTupleData size +
IndexAttributeBitMapData size ( max num filed per index + 8 - 1 /8)
        END AS index_tuple_hdr_bm,
        /* data len: we remove null values save space using it fractionnal part
from stats */
        sum( (1-coalesce(s.null_frac, 0)) * coalesce(s.avg_width, 1024)) AS
nulldatawidth,

```

```

        max( CASE WHEN i.atttypid = 'pg_catalog.name'::regtype THEN 1 ELSE 0
END ) > 0 AS is_na
    FROM (
        SELECT ct.relname AS tblname, ct.relnamespace, ic.idxname, ic.attpos,
ic.indkey, ic.indkey[ic.attpos], ic.reltuples, ic.relpages, ic.tbloid, ic.idxoid,
ic.fillfactor,
            coalesce(a1.attnum, a2.attnum) AS attnum, coalesce(a1.attname,
a2.attname) AS attname, coalesce(a1.atttypid, a2.atttypid) AS atttypid,
            CASE WHEN a1.attnum IS NULL
            THEN ic.idxname
            ELSE ct.relname
            END AS attrelname
        FROM (
            SELECT idxname, reltuples, relpages, tbloid, idxoid, fillfactor,
indkey,
                pg_catalog.generate_series(1,indnatts) AS attpos
            FROM (
                SELECT ci.relname AS idxname, ci.reltuples, ci.relpages,
i.indrelid AS tbloid,
                    i.indexrelid AS idxoid,
                    coalesce(substring(
                        array_to_string(ci.reloptions, ' ')
                        from 'fillfactor=([0-9]+)')::smallint, 90) AS fillfactor,
                    i.indnatts,
                    pg_catalog.string_to_array(pg_catalog.textin(
                        pg_catalog.int2vectorout(i.indkey)), ' ')::int[] AS indkey
                FROM pg_catalog.pg_index i
                JOIN pg_catalog.pg_class ci ON ci.oid = i.indexrelid
                WHERE ci.relam=(SELECT oid FROM pg_am WHERE amname = 'btree')
                AND ci.relpages > 0
            ) AS idx_data
        ) AS ic
        JOIN pg_catalog.pg_class ct ON ct.oid = ic.tbloid
        LEFT JOIN pg_catalog.pg_attribute a1 ON
            ic.indkey[ic.attpos] <> 0
            AND a1.attrelid = ic.tbloid
            AND a1.attnum = ic.indkey[ic.attpos]
        LEFT JOIN pg_catalog.pg_attribute a2 ON
            ic.indkey[ic.attpos] = 0
            AND a2.attrelid = ic.idxoid
            AND a2.attnum = ic.attpos
    ) i
    JOIN pg_catalog.pg_namespace n ON n.oid = i.relnamespace

```

```

JOIN pg_catalog.pg_stats s ON s.schemaname = n.nspname
                        AND s.tablename = i.attrelname
                        AND s.attnname = i.attnname
GROUP BY 1,2,3,4,5,6,7,8,9,10,11
) AS rows_data_stats
) AS rows_hdr_pdg_stats
) AS relation_stats
ORDER BY nspname, tblname, idxname;

```

Los resultados de la consulta muestran que la tabla tiene una hinchazón de alrededor del 51 por ciento:

```

current_database | schemaname | tblname | real_size | extra_size | extra_pct |
fillfactor | bloat_size | bloat_pct | is_na
-----+-----+-----+-----+-----+-----+-----
+-----+-----+-----+-----+-----+-----+-----
apgl | public | emp | 60383232 | 30744576 | 50.91575091575091 | 100 | 30744576 |
50.91575091575091 | f

```

Estas son las estadísticas de la `pg_stat_all_tables` vista:

```

relid          | 914748
schemaname     | public
relname        | emp
seq_scan       | 5
seq_tup_read    | 1500000
idx_scan       | 0
idx_tup_fetch   | 0
n_tup_ins       | 600000
n_tup_upd       | 500000
n_tup_del       | 0
n_tup_hot_upd   | 0
n_live_tup      | 500000
n_dead_tup      | 0
n_mod_since_analyze | 0
last_vacuum     |
last_autovacuum | 2023-04-15 11:59:54.957449+00
last_analyze    |
last_autoanalyze | 2023-04-15 11:59:55.016352+00
vacuum_count     | 0
autovacuum_count | 2
analyze_count    | 0

```

```
autoanalyze_count | 3
```

Observe que autovacuum actualizó las `last_autoanalyze` columnas `last_autovacuum` y después de ejecutarse.

Ahora, insertemos algunas filas en la tabla y compruébelo `extra_size(bloat_size)`, ya que el espacio vacío también se considera hinchado.

```
apgl=> select count(*) from emp;
count | 900000
```

current_database	schemaname	tblname	real_size	extra_size	extra_pct	fillfactor	bloat_size	bloat_pct	is_na
apgl	public	emp	61349888	327680	0.5341167044999332	100	327680	0.5341167044999332	f

(1 row)

La `bloat_pct` columna de la salida indica que el espacio limpio ha sido ocupado por nuevas inserciones. Vamos a correr `VACUUM FULL`:

```
apgl=> vacuum full emp ;
VACUUM
```

current_database	schemaname	tblname	real_size	extra_size	extra_pct	fillfactor	bloat_size	bloat_pct	is_na
apgl	public	emp	60792832	-229376	0	100	0	0	f

(1 row)

En este resultado, puede ver que se han eliminado el espacio vacío y la hinchazón y que el espacio se ha devuelto al sistema operativo.

Note

En lugar de hacerlo `VACUUM FULL`, puede correr `pg_repack` para obtener los mismos resultados.

Recursos

- [Descripción del autovacuum en entornos Amazon RDS para](#) [AWS PostgreSQL](#) (entrada del blog)
- [Aspiración automática \(documentación de PostgreSQL\)](#)
- [Asignación de memoria para la aspiradora automática](#) (documentación de Amazon RDS)
- [Prevención de errores envolventes en los identificadores de transacciones \(documentación de PostgreSQL\)](#)
- [Elimine la hinchazón de Amazon Aurora y RDS para PostgreSQL con pg_repack \(entrada del blog\)](#) [AWS](#)
- [Ajuste de los parámetros de PostgreSQL en Amazon RDS y Amazon AWS Aurora](#) (guía prescriptiva)

Historial de documentos

En la siguiente tabla, se describen cambios significativos de esta guía. Si quiere recibir notificaciones de futuras actualizaciones, puede suscribirse a las [notificaciones RSS](#).

Cambio	Descripción	Fecha
Se corrigió reindex la sintaxis	En la sección dedicada a la reconstrucción simultánea de un índice , se corrigieron los reindex ejemplos.	30 de junio de 2025
Publicación inicial	—	22 de diciembre de 2023

AWS Glosario de orientación prescriptiva

Los siguientes son términos de uso común en las estrategias, guías y patrones proporcionados por la Guía AWS prescriptiva. Para sugerir entradas, utilice el enlace [Enviar comentarios](#) al final del glosario.

Números

Las 7 R

Siete estrategias de migración comunes para trasladar aplicaciones a la nube. Estas estrategias se basan en las 5 R que Gartner identificó en 2011 y consisten en lo siguiente:

- **Refactorizar/rediseñar:** traslade una aplicación y modifique su arquitectura mediante el máximo aprovechamiento de las características nativas en la nube para mejorar la agilidad, el rendimiento y la escalabilidad. Por lo general, esto implica trasladar el sistema operativo y la base de datos. Ejemplo: migre su base de datos Oracle local a la edición compatible con PostgreSQL de Amazon Aurora.
- **Redefinir la plataforma (transportar y redefinir):** traslade una aplicación a la nube e introduzca algún nivel de optimización para aprovechar las capacidades de la nube. Ejemplo: migre su base de datos Oracle local a Amazon Relational Database Service (Amazon RDS) para Oracle en el. Nube de AWS
- **Recomprar (readquirir):** cambie a un producto diferente, lo cual se suele llevar a cabo al pasar de una licencia tradicional a un modelo SaaS. Ejemplo: migre su sistema de gestión de relaciones con los clientes (CRM) a Salesforce.com.
- **Volver a alojar (migrar mediante lift-and-shift):** traslade una aplicación a la nube sin realizar cambios para aprovechar las capacidades de la nube. Ejemplo: migre su base de datos Oracle local a Oracle en una EC2 instancia del. Nube de AWS
- **Reubicar:** (migrar el hipervisor mediante lift and shift): traslade la infraestructura a la nube sin comprar equipo nuevo, reescribir aplicaciones o modificar las operaciones actuales. Los servidores se migran de una plataforma local a un servicio en la nube para la misma plataforma. Ejemplo: migrar una Microsoft Hyper-V aplicación a AWS.
- **Retener (revisitar):** conserve las aplicaciones en el entorno de origen. Estas pueden incluir las aplicaciones que requieren una refactorización importante, que desee posponer para más adelante, y las aplicaciones heredadas que desee retener, ya que no hay ninguna justificación empresarial para migrarlas.

- Retirar: retire o elimine las aplicaciones que ya no sean necesarias en un entorno de origen.

A

ABAC

Consulte control de [acceso basado en atributos](#).

servicios abstractos

Consulte [servicios gestionados](#).

ACID

Consulte [atomicidad, consistencia, aislamiento y durabilidad](#).

migración activa-activa

Método de migración de bases de datos en el que las bases de datos de origen y destino se mantienen sincronizadas (mediante una herramienta de replicación bidireccional o mediante operaciones de escritura doble) y ambas bases de datos gestionan las transacciones de las aplicaciones conectadas durante la migración. Este método permite la migración en lotes pequeños y controlados, en lugar de requerir una transición única. Es más flexible, pero requiere más trabajo que la migración [activa-pasiva](#).

migración activa-pasiva

Método de migración de bases de datos en el que las bases de datos de origen y destino se mantienen sincronizadas, pero solo la base de datos de origen gestiona las transacciones de las aplicaciones conectadas, mientras los datos se replican en la base de datos de destino. La base de datos de destino no acepta ninguna transacción durante la migración.

función agregada

Función SQL que opera en un grupo de filas y calcula un único valor de retorno para el grupo. Algunos ejemplos de funciones agregadas incluyen SUM y MAX.

IA

Véase [inteligencia artificial](#).

AIOps

Consulte las [operaciones de inteligencia artificial](#).

anonimización

El proceso de eliminar permanentemente la información personal de un conjunto de datos. La anonimización puede ayudar a proteger la privacidad personal. Los datos anonimizados ya no se consideran datos personales.

antipatronos

Una solución que se utiliza con frecuencia para un problema recurrente en el que la solución es contraproducente, ineficaz o menos eficaz que una alternativa.

control de aplicaciones

Un enfoque de seguridad que permite el uso únicamente de aplicaciones aprobadas para ayudar a proteger un sistema contra el malware.

cartera de aplicaciones

Recopilación de información detallada sobre cada aplicación que utiliza una organización, incluido el costo de creación y mantenimiento de la aplicación y su valor empresarial. Esta información es clave para [el proceso de detección y análisis de la cartera](#) y ayuda a identificar y priorizar las aplicaciones que se van a migrar, modernizar y optimizar.

inteligencia artificial (IA)

El campo de la informática que se dedica al uso de tecnologías informáticas para realizar funciones cognitivas que suelen estar asociadas a los seres humanos, como el aprendizaje, la resolución de problemas y el reconocimiento de patrones. Para más información, consulte [¿Qué es la inteligencia artificial?](#)

operaciones de inteligencia artificial (AIOps)

El proceso de utilizar técnicas de machine learning para resolver problemas operativos, reducir los incidentes operativos y la intervención humana, y mejorar la calidad del servicio. Para obtener más información sobre cómo AIOps se utiliza en la estrategia de AWS migración, consulte la [guía de integración de operaciones](#).

cifrado asimétrico

Algoritmo de cifrado que utiliza un par de claves, una clave pública para el cifrado y una clave privada para el descifrado. Puede compartir la clave pública porque no se utiliza para el descifrado, pero el acceso a la clave privada debe estar sumamente restringido.

atomicidad, consistencia, aislamiento, durabilidad (ACID)

Conjunto de propiedades de software que garantizan la validez de los datos y la fiabilidad operativa de una base de datos, incluso en caso de errores, cortes de energía u otros problemas.

control de acceso basado en atributos (ABAC)

La práctica de crear permisos detallados basados en los atributos del usuario, como el departamento, el puesto de trabajo y el nombre del equipo. Para obtener más información, consulte [ABAC AWS en la](#) documentación AWS Identity and Access Management (IAM).

origen de datos fidedigno

Ubicación en la que se almacena la versión principal de los datos, que se considera la fuente de información más fiable. Puede copiar los datos del origen de datos autorizado a otras ubicaciones con el fin de procesarlos o modificarlos, por ejemplo, anonimizarlos, redactarlos o seudonimizarlos.

Zona de disponibilidad

Una ubicación distinta dentro de una Región de AWS que está aislada de los fallos en otras zonas de disponibilidad y que proporciona una conectividad de red económica y de baja latencia a otras zonas de disponibilidad de la misma región.

AWS Marco de adopción de la nube (AWS CAF)

Un marco de directrices y mejores prácticas AWS para ayudar a las organizaciones a desarrollar un plan eficiente y eficaz para migrar con éxito a la nube. AWS CAF organiza la orientación en seis áreas de enfoque denominadas perspectivas: negocios, personas, gobierno, plataforma, seguridad y operaciones. Las perspectivas empresariales, humanas y de gobernanza se centran en las habilidades y los procesos empresariales; las perspectivas de plataforma, seguridad y operaciones se centran en las habilidades y los procesos técnicos. Por ejemplo, la perspectiva humana se dirige a las partes interesadas que se ocupan de los Recursos Humanos (RR. HH.), las funciones del personal y la administración de las personas. Desde esta perspectiva, AWS CAF proporciona orientación para el desarrollo, la formación y la comunicación de las personas a fin de preparar a la organización para una adopción exitosa de la nube. Para obtener más información, consulte la [Página web de AWS CAF](#) y el [Documento técnico de AWS CAF](#).

AWS Marco de calificación de la carga de trabajo (AWS WQF)

Herramienta que evalúa las cargas de trabajo de migración de bases de datos, recomienda estrategias de migración y proporciona estimaciones de trabajo. AWS WQF se incluye con AWS

Schema Conversion Tool ().AWS SCT Analiza los esquemas de bases de datos y los objetos de código, el código de las aplicaciones, las dependencias y las características de rendimiento y proporciona informes de evaluación.

B

Un bot malo

Un [bot](#) destinado a interrumpir o causar daño a personas u organizaciones.

BCP

Consulte la [planificación de la continuidad del negocio](#).

gráfico de comportamiento

Una vista unificada e interactiva del comportamiento de los recursos y de las interacciones a lo largo del tiempo. Puede utilizar un gráfico de comportamiento con Amazon Detective para examinar los intentos de inicio de sesión fallidos, las llamadas sospechosas a la API y acciones similares. Para obtener más información, consulte [Datos en un gráfico de comportamiento](#) en la documentación de Detective.

sistema big-endian

Un sistema que almacena primero el byte más significativo. Véase también [endianness](#).

clasificación binaria

Un proceso que predice un resultado binario (una de las dos clases posibles). Por ejemplo, es posible que su modelo de ML necesite predecir problemas como “¿Este correo electrónico es spam o no es spam?” o “¿Este producto es un libro o un automóvil?”.

filtro de floración

Estructura de datos probabilística y eficiente en términos de memoria que se utiliza para comprobar si un elemento es miembro de un conjunto.

implementación azul/verde

Una estrategia de despliegue en la que se crean dos entornos separados pero idénticos. La versión actual de la aplicación se ejecuta en un entorno (azul) y la nueva versión de la aplicación en el otro entorno (verde). Esta estrategia le ayuda a revertirla rápidamente con un impacto mínimo.

bot

Aplicación de software que ejecuta tareas automatizadas a través de Internet y simula la actividad o interacción humana. Algunos bots son útiles o beneficiosos, como los rastreadores web que indexan información en Internet. Algunos otros bots, conocidos como bots malos, tienen como objetivo interrumpir o causar daños a personas u organizaciones.

botnet

Redes de [bots](#) que están infectadas por [malware](#) y que están bajo el control de una sola parte, conocida como pastor u operador de bots. Las botnets son el mecanismo más conocido para escalar los bots y su impacto.

branch

Área contenida de un repositorio de código. La primera rama que se crea en un repositorio es la rama principal. Puede crear una rama nueva a partir de una rama existente y, a continuación, desarrollar características o corregir errores en la rama nueva. Una rama que se genera para crear una característica se denomina comúnmente rama de característica. Cuando la característica se encuentra lista para su lanzamiento, se vuelve a combinar la rama de característica con la rama principal. Para obtener más información, consulte [Acerca de las sucursales](#) (GitHub documentación).

acceso con cristales rotos

En circunstancias excepcionales y mediante un proceso aprobado, un usuario puede acceder rápidamente a un sitio para el Cuenta de AWS que normalmente no tiene permisos de acceso. Para obtener más información, consulte el indicador [Implemente procedimientos de rotura de cristales en la guía Well-Architected](#) AWS .

estrategia de implementación sobre infraestructura existente

La infraestructura existente en su entorno. Al adoptar una estrategia de implementación sobre infraestructura existente para una arquitectura de sistemas, se diseña la arquitectura en función de las limitaciones de los sistemas y la infraestructura actuales. Si está ampliando la infraestructura existente, puede combinar las estrategias de implementación sobre infraestructuras existentes y de [implementación desde cero](#).

caché de búfer

El área de memoria donde se almacenan los datos a los que se accede con más frecuencia.

capacidad empresarial

Lo que hace una empresa para generar valor (por ejemplo, ventas, servicio al cliente o marketing). Las arquitecturas de microservicios y las decisiones de desarrollo pueden estar impulsadas por las capacidades empresariales. Para obtener más información, consulte la sección [Organizado en torno a las capacidades empresariales](#) del documento técnico [Ejecutar microservicios en contenedores en AWS](#).

planificación de la continuidad del negocio (BCP)

Plan que aborda el posible impacto de un evento disruptivo, como una migración a gran escala en las operaciones y permite a la empresa reanudar las operaciones rápidamente.

C

CAF

[Consulte el marco AWS de adopción de la nube.](#)

despliegue canario

El lanzamiento lento e incremental de una versión para los usuarios finales. Cuando está seguro, despliega la nueva versión y reemplaza la versión actual en su totalidad.

CCoE

Consulte [Cloud Center of Excellence](#).

CDC

Consulte la [captura de datos de cambios](#).

captura de datos de cambio (CDC)

Proceso de seguimiento de los cambios en un origen de datos, como una tabla de base de datos, y registro de los metadatos relacionados con el cambio. Puede utilizar los CDC para diversos fines, como auditar o replicar los cambios en un sistema de destino para mantener la sincronización.

ingeniería del caos

Introducir intencionalmente fallos o eventos disruptivos para poner a prueba la resiliencia de un sistema. Puedes usar [AWS Fault Injection Service \(AWS FIS\)](#) para realizar experimentos que estresen tus AWS cargas de trabajo y evalúen su respuesta.

CI/CD

Consulte la [integración continua y la entrega continua](#).

clasificación

Un proceso de categorización que permite generar predicciones. Los modelos de ML para problemas de clasificación predicen un valor discreto. Los valores discretos siempre son distintos entre sí. Por ejemplo, es posible que un modelo necesite evaluar si hay o no un automóvil en una imagen.

cifrado del cliente

Cifrado de datos localmente, antes de que el objetivo los Servicio de AWS reciba.

Centro de excelencia en la nube (CCoE)

Equipo multidisciplinario que impulsa los esfuerzos de adopción de la nube en toda la organización, incluido el desarrollo de las prácticas recomendadas en la nube, la movilización de recursos, el establecimiento de plazos de migración y la dirección de la organización durante las transformaciones a gran escala. Para obtener más información, consulte las [publicaciones de CCo E](#) en el blog de estrategia Nube de AWS empresarial.

computación en la nube

La tecnología en la nube que se utiliza normalmente para la administración de dispositivos de IoT y el almacenamiento de datos de forma remota. La computación en la nube suele estar conectada a la tecnología de [computación perimetral](#).

modelo operativo en la nube

En una organización de TI, el modelo operativo que se utiliza para crear, madurar y optimizar uno o más entornos de nube. Para obtener más información, consulte [Creación de su modelo operativo de nube](#).

etapas de adopción de la nube

Las cuatro fases por las que suelen pasar las organizaciones cuando migran a Nube de AWS:

- Proyecto: ejecución de algunos proyectos relacionados con la nube con fines de prueba de concepto y aprendizaje
- Fundamento: realizar inversiones fundamentales para escalar su adopción de la nube (p. ej., crear una landing zone, definir una CCo E, establecer un modelo de operaciones)

- Migración: migración de aplicaciones individuales
- Reinención: optimización de productos y servicios e innovación en la nube

Stephen Orban definió estas etapas en la entrada del blog [The Journey Toward Cloud-First & the Stages of Adoption en el](#) blog Nube de AWS Enterprise Strategy. Para obtener información sobre su relación con la estrategia de AWS migración, consulte la guía de [preparación para la migración](#).

CMDB

Consulte la [base de datos de administración de la configuración](#).

repositorio de código

Una ubicación donde el código fuente y otros activos, como documentación, muestras y scripts, se almacenan y actualizan mediante procesos de control de versiones. Los repositorios en la nube más comunes incluyen GitHub o Bitbucket Cloud. Cada versión del código se denomina rama. En una estructura de microservicios, cada repositorio se encuentra dedicado a una única funcionalidad. Una sola canalización de CI/CD puede utilizar varios repositorios.

caché en frío

Una caché de búfer que está vacía no está bien poblada o contiene datos obsoletos o irrelevantes. Esto afecta al rendimiento, ya que la instancia de la base de datos debe leer desde la memoria principal o el disco, lo que es más lento que leer desde la memoria caché del búfer.

datos fríos

Datos a los que se accede con poca frecuencia y que suelen ser históricos. Al consultar este tipo de datos, normalmente se aceptan consultas lentas. Trasladar estos datos a niveles o clases de almacenamiento de menor rendimiento y menos costosos puede reducir los costos.

visión artificial (CV)

Campo de la [IA](#) que utiliza el aprendizaje automático para analizar y extraer información de formatos visuales, como imágenes y vídeos digitales. Por ejemplo, Amazon SageMaker AI proporciona algoritmos de procesamiento de imágenes para CV.

desviación de configuración

En el caso de una carga de trabajo, un cambio de configuración con respecto al estado esperado. Puede provocar que la carga de trabajo deje de cumplir las normas y, por lo general, es gradual e involuntario.

base de datos de administración de configuración (CMDB)

Repositorio que almacena y administra información sobre una base de datos y su entorno de TI, incluidos los componentes de hardware y software y sus configuraciones. Por lo general, los datos de una CMDB se utilizan en la etapa de detección y análisis de la cartera de productos durante la migración.

paquete de conformidad

Conjunto de AWS Config reglas y medidas correctivas que puede reunir para personalizar sus comprobaciones de conformidad y seguridad. Puede implementar un paquete de conformidad como una entidad única en una región Cuenta de AWS y, o en una organización, mediante una plantilla YAML. Para obtener más información, consulta los [paquetes de conformidad](#) en la documentación. AWS Config

integración y entrega continuas (CI/CD)

El proceso de automatización de las etapas de origen, compilación, prueba, puesta en escena y producción del proceso de publicación del software. CI/CD is commonly described as a pipeline. CI/CD puede ayudarlo a automatizar los procesos, mejorar la productividad, mejorar la calidad del código y entregar con mayor rapidez. Para obtener más información, consulte [Beneficios de la entrega continua](#). CD también puede significar implementación continua. Para obtener más información, consulte [Entrega continua frente a implementación continua](#).

CV

Vea la [visión artificial](#).

D

datos en reposo

Datos que están estacionarios en la red, como los datos que se encuentran almacenados.

clasificación de datos

Un proceso para identificar y clasificar los datos de su red en función de su importancia y sensibilidad. Es un componente fundamental de cualquier estrategia de administración de riesgos de ciberseguridad porque lo ayuda a determinar los controles de protección y retención adecuados para los datos. La clasificación de datos es un componente del pilar de seguridad

del AWS Well-Architected Framework. Para obtener más información, consulte [Clasificación de datos](#).

desviación de datos

Una variación significativa entre los datos de producción y los datos que se utilizaron para entrenar un modelo de machine learning, o un cambio significativo en los datos de entrada a lo largo del tiempo. La desviación de los datos puede reducir la calidad, la precisión y la imparcialidad generales de las predicciones de los modelos de machine learning.

datos en tránsito

Datos que se mueven de forma activa por la red, por ejemplo, entre los recursos de la red.

mallado de datos

Un marco arquitectónico que proporciona una propiedad de datos distribuida y descentralizada con una administración y un gobierno centralizados.

minimización de datos

El principio de recopilar y procesar solo los datos estrictamente necesarios. Practicar la minimización de los datos Nube de AWS puede reducir los riesgos de privacidad, los costos y la huella de carbono de la analítica.

perímetro de datos

Un conjunto de barreras preventivas en su AWS entorno que ayudan a garantizar que solo las identidades confiables accedan a los recursos confiables desde las redes esperadas. Para obtener más información, consulte [Crear un perímetro de datos sobre](#) AWS

preprocesamiento de datos

Transformar los datos sin procesar en un formato que su modelo de ML pueda analizar fácilmente. El preprocesamiento de datos puede implicar eliminar determinadas columnas o filas y corregir los valores faltantes, incoherentes o duplicados.

procedencia de los datos

El proceso de rastrear el origen y el historial de los datos a lo largo de su ciclo de vida, por ejemplo, la forma en que se generaron, transmitieron y almacenaron los datos.

titular de los datos

Persona cuyos datos se recopilan y procesan.

almacenamiento de datos

Un sistema de administración de datos que respalde la inteligencia empresarial, como la analítica. Los almacenes de datos suelen contener grandes cantidades de datos históricos y, por lo general, se utilizan para consultas y análisis.

lenguaje de definición de datos (DDL)

Instrucciones o comandos para crear o modificar la estructura de tablas y objetos de una base de datos.

lenguaje de manipulación de datos (DML)

Instrucciones o comandos para modificar (insertar, actualizar y eliminar) la información de una base de datos.

DDL

Consulte el [lenguaje de definición de bases](#) de datos.

conjunto profundo

Combinar varios modelos de aprendizaje profundo para la predicción. Puede utilizar conjuntos profundos para obtener una predicción más precisa o para estimar la incertidumbre de las predicciones.

aprendizaje profundo

Un subcampo del ML que utiliza múltiples capas de redes neuronales artificiales para identificar el mapeo entre los datos de entrada y las variables objetivo de interés.

defense-in-depth

Un enfoque de seguridad de la información en el que se distribuyen cuidadosamente una serie de mecanismos y controles de seguridad en una red informática para proteger la confidencialidad, la integridad y la disponibilidad de la red y de los datos que contiene. Al adoptar esta estrategia AWS, se añaden varios controles en diferentes capas de la AWS Organizations estructura para ayudar a proteger los recursos. Por ejemplo, un defense-in-depth enfoque podría combinar la autenticación multifactorial, la segmentación de la red y el cifrado.

administrador delegado

En AWS Organizations, un servicio compatible puede registrar una cuenta de AWS miembro para administrar las cuentas de la organización y gestionar los permisos de ese servicio. Esta

cuenta se denomina administrador delegado para ese servicio. Para obtener más información y una lista de servicios compatibles, consulte [Servicios que funcionan con AWS Organizations](#) en la documentación de AWS Organizations .

Implementación

El proceso de hacer que una aplicación, características nuevas o correcciones de código se encuentren disponibles en el entorno de destino. La implementación abarca implementar cambios en una base de código y, a continuación, crear y ejecutar esa base en los entornos de la aplicación.

entorno de desarrollo

Consulte [entorno](#).

control de detección

Un control de seguridad que se ha diseñado para detectar, registrar y alertar después de que se produzca un evento. Estos controles son una segunda línea de defensa, ya que lo advierten sobre los eventos de seguridad que han eludido los controles preventivos establecidos. Para obtener más información, consulte [Controles de detección](#) en Implementación de controles de seguridad en AWS.

asignación de flujos de valor para el desarrollo (DVSM)

Proceso que se utiliza para identificar y priorizar las restricciones que afectan negativamente a la velocidad y la calidad en el ciclo de vida del desarrollo de software. DVSM amplía el proceso de asignación del flujo de valor diseñado originalmente para las prácticas de fabricación ajustada. Se centra en los pasos y los equipos necesarios para crear y transferir valor a través del proceso de desarrollo de software.

gemelo digital

Representación virtual de un sistema del mundo real, como un edificio, una fábrica, un equipo industrial o una línea de producción. Los gemelos digitales son compatibles con el mantenimiento predictivo, la supervisión remota y la optimización de la producción.

tabla de dimensiones

En un [esquema en estrella](#), tabla más pequeña que contiene los atributos de datos sobre los datos cuantitativos de una tabla de hechos. Los atributos de la tabla de dimensiones suelen ser campos de texto o números discretos que se comportan como texto. Estos atributos se utilizan habitualmente para restringir consultas, filtrar y etiquetar conjuntos de resultados.

desastre

Un evento que impide que una carga de trabajo o un sistema cumplan sus objetivos empresariales en su ubicación principal de implementación. Estos eventos pueden ser desastres naturales, fallos técnicos o el resultado de acciones humanas, como una configuración incorrecta involuntaria o un ataque de malware.

recuperación de desastres (DR)

La estrategia y el proceso que se utilizan para minimizar el tiempo de inactividad y la pérdida de datos ocasionados por un [desastre](#). Para obtener más información, consulte [Recuperación ante desastres de cargas de trabajo en AWS: Recovery in the Cloud in the AWS Well-Architected Framework](#).

DML

Consulte el lenguaje de manipulación de [bases de datos](#).

diseño basado en el dominio

Un enfoque para desarrollar un sistema de software complejo mediante la conexión de sus componentes a dominios en evolución, o a los objetivos empresariales principales, a los que sirve cada componente. Este concepto lo introdujo Eric Evans en su libro, *Diseño impulsado por el dominio: abordando la complejidad en el corazón del software* (Boston: Addison-Wesley Professional, 2003). Para obtener información sobre cómo utilizar el diseño basado en dominios con el patrón de higos estranguladores, consulte [Modernización gradual de los servicios web antiguos de Microsoft ASP.NET \(ASMX\) mediante contenedores y Amazon API Gateway](#).

DR

Consulte [recuperación ante desastres](#).

detección de deriva

Seguimiento de las desviaciones con respecto a una configuración de referencia. Por ejemplo, puedes usarlo AWS CloudFormation para [detectar desviaciones en los recursos del sistema](#) o puedes usarlo AWS Control Tower para [detectar cambios en tu landing zone](#) que puedan afectar al cumplimiento de los requisitos de gobierno.

DVSM

Consulte [el mapeo del flujo de valor del desarrollo](#).

E

EDA

Consulte el [análisis exploratorio de datos](#).

EDI

Véase [intercambio electrónico de datos](#).

computación en la periferia

La tecnología que aumenta la potencia de cálculo de los dispositivos inteligentes en la periferia de una red de IoT. En comparación con [la computación en nube](#), [la computación](#) perimetral puede reducir la latencia de la comunicación y mejorar el tiempo de respuesta.

intercambio electrónico de datos (EDI)

El intercambio automatizado de documentos comerciales entre organizaciones. Para obtener más información, consulte [Qué es el intercambio electrónico de datos](#).

cifrado

Proceso informático que transforma datos de texto plano, legibles por humanos, en texto cifrado.

clave de cifrado

Cadena criptográfica de bits aleatorios que se genera mediante un algoritmo de cifrado. Las claves pueden variar en longitud y cada una se ha diseñado para ser impredecible y única.

endianidad

El orden en el que se almacenan los bytes en la memoria del ordenador. Los sistemas big-endianos almacenan primero el byte más significativo. Los sistemas Little-Endian almacenan primero el byte menos significativo.

punto de conexión

[Consulte el punto final del servicio](#).

servicio de punto de conexión

Servicio que puede alojar en una nube privada virtual (VPC) para compartir con otros usuarios. Puede crear un servicio de punto final AWS PrivateLink y conceder permisos a otros directores

Cuentas de AWS o a AWS Identity and Access Management (IAM). Estas cuentas o entidades principales pueden conectarse a su servicio de punto de conexión de forma privada mediante la creación de puntos de conexión de VPC de interfaz. Para obtener más información, consulte [Creación de un servicio de punto de conexión](#) en la documentación de Amazon Virtual Private Cloud (Amazon VPC).

planificación de recursos empresariales (ERP)

Un sistema que automatiza y gestiona los procesos empresariales clave (como la contabilidad, el [MES](#) y la gestión de proyectos) de una empresa.

cifrado de sobre

El proceso de cifrar una clave de cifrado con otra clave de cifrado. Para obtener más información, consulte el [cifrado de sobres](#) en la documentación de AWS Key Management Service (AWS KMS).

entorno

Una instancia de una aplicación en ejecución. Los siguientes son los tipos de entornos más comunes en la computación en la nube:

- entorno de desarrollo: instancia de una aplicación en ejecución que solo se encuentra disponible para el equipo principal responsable del mantenimiento de la aplicación. Los entornos de desarrollo se utilizan para probar los cambios antes de promocionarlos a los entornos superiores. Este tipo de entorno a veces se denomina entorno de prueba.
- entornos inferiores: todos los entornos de desarrollo de una aplicación, como los que se utilizan para las compilaciones y pruebas iniciales.
- entorno de producción: instancia de una aplicación en ejecución a la que pueden acceder los usuarios finales. En una canalización de CI/CD, el entorno de producción es el último entorno de implementación.
- entornos superiores: todos los entornos a los que pueden acceder usuarios que no sean del equipo de desarrollo principal. Esto puede incluir un entorno de producción, entornos de preproducción y entornos para las pruebas de aceptación por parte de los usuarios.

epopeya

En las metodologías ágiles, son categorías funcionales que ayudan a organizar y priorizar el trabajo. Las epopeyas brindan una descripción detallada de los requisitos y las tareas de implementación. Por ejemplo, las epopeyas AWS de seguridad de CAF incluyen la gestión de identidades y accesos, los controles de detección, la seguridad de la infraestructura, la protección

de datos y la respuesta a incidentes. Para obtener más información sobre las epopeyas en la estrategia de migración de AWS , consulte la [Guía de implementación del programa](#).

ERP

Consulte [planificación de recursos empresariales](#).

análisis de datos de tipo exploratorio (EDA)

El proceso de analizar un conjunto de datos para comprender sus características principales. Se recopilan o agregan datos y, a continuación, se realizan las investigaciones iniciales para encontrar patrones, detectar anomalías y comprobar las suposiciones. El EDA se realiza mediante el cálculo de estadísticas resumidas y la creación de visualizaciones de datos.

F

tabla de datos

La tabla central de un [esquema en forma de estrella](#). Almacena datos cuantitativos sobre las operaciones comerciales. Normalmente, una tabla de hechos contiene dos tipos de columnas: las que contienen medidas y las que contienen una clave externa para una tabla de dimensiones.

fallan rápidamente

Una filosofía que utiliza pruebas frecuentes e incrementales para reducir el ciclo de vida del desarrollo. Es una parte fundamental de un enfoque ágil.

límite de aislamiento de fallas

En el Nube de AWS, un límite, como una zona de disponibilidad Región de AWS, un plano de control o un plano de datos, que limita el efecto de una falla y ayuda a mejorar la resiliencia de las cargas de trabajo. Para obtener más información, consulte [Límites de AWS aislamiento](#) de errores.

rama de característica

Consulte la [sucursal](#).

características

Los datos de entrada que se utilizan para hacer una predicción. Por ejemplo, en un contexto de fabricación, las características pueden ser imágenes que se capturan periódicamente desde la línea de fabricación.

importancia de las características

La importancia que tiene una característica para las predicciones de un modelo. Por lo general, esto se expresa como una puntuación numérica que se puede calcular mediante diversas técnicas, como las explicaciones aditivas de Shapley (SHAP) y los gradientes integrados. Para obtener más información, consulte [Interpretabilidad del modelo de aprendizaje automático con AWS](#).

transformación de funciones

Optimizar los datos para el proceso de ML, lo que incluye enriquecer los datos con fuentes adicionales, escalar los valores o extraer varios conjuntos de información de un solo campo de datos. Esto permite que el modelo de ML se beneficie de los datos. Por ejemplo, si divide la fecha del “27 de mayo de 2021 00:15:37” en “jueves”, “mayo”, “2021” y “15”, puede ayudar al algoritmo de aprendizaje a aprender patrones matizados asociados a los diferentes componentes de los datos.

indicaciones de unos pocos pasos

Proporcionar a un [LLM](#) un pequeño número de ejemplos que demuestren la tarea y el resultado deseado antes de pedirle que realice una tarea similar. Esta técnica es una aplicación del aprendizaje contextual, en el que los modelos aprenden a partir de ejemplos (planos) integrados en las instrucciones. Las indicaciones con pocas tomas pueden ser eficaces para tareas que requieren un formato, un razonamiento o un conocimiento del dominio específicos. [Consulte también el apartado de mensajes sin intervención](#).

FGAC

Consulte el control [de acceso detallado](#).

control de acceso preciso (FGAC)

El uso de varias condiciones que tienen por objetivo permitir o denegar una solicitud de acceso.

migración relámpago

Método de migración de bases de datos que utiliza la replicación continua de datos mediante la [captura de datos modificados](#) para migrar los datos en el menor tiempo posible, en lugar de utilizar un enfoque gradual. El objetivo es reducir al mínimo el tiempo de inactividad.

FM

Consulte el [modelo básico](#).

modelo de base (FM)

Una gran red neuronal de aprendizaje profundo que se ha estado entrenando con conjuntos de datos masivos de datos generalizados y sin etiquetar. FMs son capaces de realizar una amplia variedad de tareas generales, como comprender el lenguaje, generar texto e imágenes y conversar en lenguaje natural. Para obtener más información, consulte [Qué son los modelos básicos](#).

G

IA generativa

Un subconjunto de modelos de [IA](#) que se han entrenado con grandes cantidades de datos y que pueden utilizar un simple mensaje de texto para crear contenido y artefactos nuevos, como imágenes, vídeos, texto y audio. Para obtener más información, consulte [Qué es la IA generativa](#).

bloqueo geográfico

Consulta [las restricciones geográficas](#).

restricciones geográficas (bloqueo geográfico)

En Amazon CloudFront, una opción para impedir que los usuarios de países específicos accedan a las distribuciones de contenido. Puede utilizar una lista de permitidos o bloqueados para especificar los países aprobados y prohibidos. Para obtener más información, consulta [Restringir la distribución geográfica del contenido](#) en la CloudFront documentación.

Flujo de trabajo de Gitflow

Un enfoque en el que los entornos inferiores y superiores utilizan diferentes ramas en un repositorio de código fuente. El flujo de trabajo de Gitflow se considera heredado, y el [flujo de trabajo basado en enlaces troncales](#) es el enfoque moderno preferido.

imagen dorada

Instantánea de un sistema o software que se utiliza como plantilla para implementar nuevas instancias de ese sistema o software. Por ejemplo, en la fabricación, una imagen dorada se puede utilizar para aprovisionar software en varios dispositivos y ayuda a mejorar la velocidad, la escalabilidad y la productividad de las operaciones de fabricación de dispositivos.

estrategia de implementación desde cero

La ausencia de infraestructura existente en un entorno nuevo. Al adoptar una estrategia de implementación desde cero para una arquitectura de sistemas, puede seleccionar todas las tecnologías nuevas sin que estas deban ser compatibles con una infraestructura existente, lo que también se conoce como [implementación sobre infraestructura existente](#). Si está ampliando la infraestructura existente, puede combinar las estrategias de implementación sobre infraestructuras existentes y de implementación desde cero.

barrera de protección

Una regla de alto nivel que ayuda a regular los recursos, las políticas y el cumplimiento en todas las unidades organizativas (OUs). Las barreras de protección preventivas aplican políticas para garantizar la alineación con los estándares de conformidad. Se implementan mediante políticas de control de servicios y límites de permisos de IAM. Las barreras de protección de detección detectan las vulneraciones de las políticas y los problemas de conformidad, y generan alertas para su corrección. Se implementan mediante Amazon AWS Config AWS Security Hub GuardDuty AWS Trusted Advisor, Amazon Inspector y AWS Lambda cheques personalizados.

H

HA

Consulte la [alta disponibilidad](#).

migración heterogénea de bases de datos

Migración de la base de datos de origen a una base de datos de destino que utilice un motor de base de datos diferente (por ejemplo, de Oracle a Amazon Aurora). La migración heterogénea suele ser parte de un esfuerzo de rediseño de la arquitectura y convertir el esquema puede ser una tarea compleja. [AWS ofrece AWS SCT](#), lo cual ayuda con las conversiones de esquemas.

alta disponibilidad (HA)

La capacidad de una carga de trabajo para funcionar de forma continua, sin intervención, en caso de desafíos o desastres. Los sistemas de alta disponibilidad están diseñados para realizar una conmutación por error automática, ofrecer un rendimiento de alta calidad de forma constante y gestionar diferentes cargas y fallos con un impacto mínimo en el rendimiento.

modernización histórica

Un enfoque utilizado para modernizar y actualizar los sistemas de tecnología operativa (TO) a fin de satisfacer mejor las necesidades de la industria manufacturera. Un histórico es un tipo de base de datos que se utiliza para recopilar y almacenar datos de diversas fuentes en una fábrica.

datos retenidos

Parte de los datos históricos etiquetados que se ocultan de un conjunto de datos que se utiliza para entrenar un modelo de aprendizaje [automático](#). Puede utilizar los datos de reserva para evaluar el rendimiento del modelo comparando las predicciones del modelo con los datos de reserva.

migración homogénea de bases de datos

Migración de la base de datos de origen a una base de datos de destino que comparte el mismo motor de base de datos (por ejemplo, Microsoft SQL Server a Amazon RDS para SQL Server). La migración homogénea suele formar parte de un esfuerzo para volver a alojar o redefinir la plataforma. Puede utilizar las utilidades de bases de datos nativas para migrar el esquema.

datos recientes

Datos a los que se accede con frecuencia, como datos en tiempo real o datos traslacionales recientes. Por lo general, estos datos requieren un nivel o una clase de almacenamiento de alto rendimiento para proporcionar respuestas rápidas a las consultas.

hotfix

Una solución urgente para un problema crítico en un entorno de producción. Debido a su urgencia, las revisiones suelen realizarse fuera del flujo de trabajo habitual de las versiones.

DevOps

periodo de hiperatención

Periodo, inmediatamente después de la transición, durante el cual un equipo de migración administra y monitorea las aplicaciones migradas en la nube para solucionar cualquier problema. Por lo general, este periodo dura de 1 a 4 días. Al final del periodo de hiperatención, el equipo de migración suele transferir la responsabilidad de las aplicaciones al equipo de operaciones en la nube.

I

IaC

Vea [la infraestructura como código](#).

políticas basadas en identidad

Política asociada a uno o más directores de IAM que define sus permisos en el Nube de AWS entorno.

aplicación inactiva

Aplicación que utiliza un promedio de CPU y memoria de entre 5 y 20 por ciento durante un periodo de 90 días. En un proyecto de migración, es habitual retirar estas aplicaciones o mantenerlas en las instalaciones.

IIoT

Consulte [Internet de las cosas industrial](#).

infraestructura inmutable

Un modelo que implementa una nueva infraestructura para las cargas de trabajo de producción en lugar de actualizar, aplicar parches o modificar la infraestructura existente. [Las infraestructuras inmutables son intrínsecamente más consistentes, fiables y predecibles que las infraestructuras mutables](#). Para obtener más información, consulte las prácticas recomendadas para [implementar con una infraestructura inmutable](#) en Well-Architected Framework AWS .

VPC entrante (de entrada)

En una arquitectura de AWS cuentas múltiples, una VPC que acepta, inspecciona y enruta las conexiones de red desde fuera de una aplicación. La [arquitectura AWS de referencia de seguridad](#) recomienda configurar la cuenta de red con entradas, salidas e inspección VPCs para proteger la interfaz bidireccional entre la aplicación y el resto de Internet.

migración gradual

Estrategia de transición en la que se migra la aplicación en partes pequeñas en lugar de realizar una transición única y completa. Por ejemplo, puede trasladar inicialmente solo unos pocos microservicios o usuarios al nuevo sistema. Tras comprobar que todo funciona correctamente, puede trasladar microservicios o usuarios adicionales de forma gradual hasta que pueda retirar su sistema heredado. Esta estrategia reduce los riesgos asociados a las grandes migraciones.

Industria 4.0

Un término que [Klaus Schwab](#) introdujo en 2016 para referirse a la modernización de los procesos de fabricación mediante avances en la conectividad, los datos en tiempo real, la automatización, el análisis y la inteligencia artificial/aprendizaje automático.

infraestructura

Todos los recursos y activos que se encuentran en el entorno de una aplicación.

infraestructura como código (IaC)

Proceso de aprovisionamiento y administración de la infraestructura de una aplicación mediante un conjunto de archivos de configuración. La IaC se ha diseñado para ayudarlo a centralizar la administración de la infraestructura, estandarizar los recursos y escalar con rapidez a fin de que los entornos nuevos sean repetibles, fiables y consistentes.

Internet de las cosas industrial (IIoT)

El uso de sensores y dispositivos conectados a Internet en los sectores industriales, como el productivo, el eléctrico, el automotriz, el sanitario, el de las ciencias de la vida y el de la agricultura. Para obtener más información, consulte [Creación de una estrategia de transformación digital de la Internet de las cosas \(IIoT\) industrial](#).

VPC de inspección

En una arquitectura de AWS cuentas múltiples, una VPC centralizada que gestiona las inspecciones del tráfico de red VPCs entre Internet y las redes locales (en una misma o Regiones de AWS diferente). La [arquitectura AWS de referencia de seguridad](#) recomienda configurar su cuenta de red con entrada, salida e inspección VPCs para proteger la interfaz bidireccional entre la aplicación e Internet en general.

Internet de las cosas (IoT)

Red de objetos físicos conectados con sensores o procesadores integrados que se comunican con otros dispositivos y sistemas a través de Internet o de una red de comunicación local. Para obtener más información, consulte [¿Qué es IoT?](#).

interpretabilidad

Característica de un modelo de machine learning que describe el grado en que un ser humano puede entender cómo las predicciones del modelo dependen de sus entradas. Para obtener más información, consulte Interpretabilidad del [modelo de aprendizaje automático](#) con AWS

IoT

Consulte [Internet de las cosas](#).

biblioteca de información de TI (ITIL)

Conjunto de prácticas recomendadas para ofrecer servicios de TI y alinearlos con los requisitos empresariales. La ITIL proporciona la base para la ITSM.

administración de servicios de TI (ITSM)

Actividades asociadas con el diseño, la implementación, la administración y el soporte de los servicios de TI para una organización. Para obtener información sobre la integración de las operaciones en la nube con las herramientas de ITSM, consulte la [Guía de integración de operaciones](#).

ITIL

Consulte la [biblioteca de información de TI](#).

ITSM

Consulte [Administración de servicios de TI](#).

L

control de acceso basado en etiquetas (LBAC)

Una implementación del control de acceso obligatorio (MAC) en la que a los usuarios y a los propios datos se les asigna explícitamente un valor de etiqueta de seguridad. La intersección entre la etiqueta de seguridad del usuario y la etiqueta de seguridad de los datos determina qué filas y columnas puede ver el usuario.

zona de aterrizaje

Una landing zone es un AWS entorno multicuenta bien diseñado, escalable y seguro. Este es un punto de partida desde el cual las empresas pueden lanzar e implementar rápidamente cargas de trabajo y aplicaciones con confianza en su entorno de seguridad e infraestructura. Para obtener más información sobre las zonas de aterrizaje, consulte [Configuración de un entorno de AWS seguro y escalable con varias cuentas](#).

modelo de lenguaje grande (LLM)

Un modelo de [IA](#) de aprendizaje profundo que se entrena previamente con una gran cantidad de datos. Un LLM puede realizar múltiples tareas, como responder preguntas, resumir documentos, traducir textos a otros idiomas y completar oraciones. [Para obtener más información, consulte Qué son. LLMs](#)

migración grande

Migración de 300 servidores o más.

LBAC

Consulte control de [acceso basado en etiquetas](#).

privilegio mínimo

La práctica recomendada de seguridad que consiste en conceder los permisos mínimos necesarios para realizar una tarea. Para obtener más información, consulte [Aplicar permisos de privilegio mínimo](#) en la documentación de IAM.

migrar mediante lift-and-shift

Ver [7 Rs](#).

sistema little-endian

Un sistema que almacena primero el byte menos significativo. Véase también [endianness](#).

LLM

Véase un modelo de lenguaje [amplio](#).

entornos inferiores

Véase [entorno](#).

M

machine learning (ML)

Un tipo de inteligencia artificial que utiliza algoritmos y técnicas para el reconocimiento y el aprendizaje de patrones. El ML analiza y aprende de los datos registrados, como los datos del

Internet de las cosas (IoT), para generar un modelo estadístico basado en patrones. Para más información, consulte [Machine learning](#).

rama principal

Ver [sucursal](#).

malware

Software diseñado para comprometer la seguridad o la privacidad de la computadora. El malware puede interrumpir los sistemas informáticos, filtrar información confidencial u obtener acceso no autorizado. Algunos ejemplos de malware son los virus, los gusanos, el ransomware, los troyanos, el spyware y los registradores de pulsaciones de teclas.

servicios gestionados

Servicios de AWS para los que AWS opera la capa de infraestructura, el sistema operativo y las plataformas, y usted accede a los puntos finales para almacenar y recuperar datos. Amazon Simple Storage Service (Amazon S3) y Amazon DynamoDB son ejemplos de servicios gestionados. También se conocen como servicios abstractos.

sistema de ejecución de fabricación (MES)

Un sistema de software para rastrear, monitorear, documentar y controlar los procesos de producción que convierten las materias primas en productos terminados en el taller.

MAP

Consulte [Migration Acceleration Program](#).

mecanismo

Un proceso completo en el que se crea una herramienta, se impulsa su adopción y, a continuación, se inspeccionan los resultados para realizar ajustes. Un mecanismo es un ciclo que se refuerza y mejora a sí mismo a medida que funciona. Para obtener más información, consulte [Creación de mecanismos](#) en el AWS Well-Architected Framework.

cuenta de miembro

Todas las Cuentas de AWS demás cuentas, excepto la de administración, que forman parte de una organización. AWS Organizations Una cuenta no puede pertenecer a más de una organización a la vez.

MES

Consulte el [sistema de ejecución de la fabricación](#).

Transporte telemétrico de Message Queue Queue (MQTT)

[Un protocolo de comunicación ligero machine-to-machine \(M2M\), basado en el patrón de publicación/suscripción, para dispositivos de IoT con recursos limitados.](#)

microservicio

Un servicio pequeño e independiente que se comunica a través de una red bien definida APIs y que, por lo general, es propiedad de equipos pequeños e independientes. Por ejemplo, un sistema de seguros puede incluir microservicios que se adapten a las capacidades empresariales, como las de ventas o marketing, o a subdominios, como las de compras, reclamaciones o análisis. Los beneficios de los microservicios incluyen la agilidad, la escalabilidad flexible, la facilidad de implementación, el código reutilizable y la resiliencia. Para obtener más información, consulte [Integrar microservicios mediante AWS servicios sin servidor](#).

arquitectura de microservicios

Un enfoque para crear una aplicación con componentes independientes que ejecutan cada proceso de la aplicación como un microservicio. Estos microservicios se comunican a través de una interfaz bien definida mediante un uso ligero. APIs Cada microservicio de esta arquitectura se puede actualizar, implementar y escalar para satisfacer la demanda de funciones específicas de una aplicación. Para obtener más información, consulte [Implementación de microservicios](#) en AWS.

Programa de aceleración de la migración (MAP)

Un AWS programa que proporciona soporte de consultoría, formación y servicios para ayudar a las organizaciones a crear una base operativa sólida para migrar a la nube y para ayudar a compensar el costo inicial de las migraciones. El MAP incluye una metodología de migración para ejecutar las migraciones antiguas de forma metódica y un conjunto de herramientas para automatizar y acelerar los escenarios de migración más comunes.

migración a escala

Proceso de transferencia de la mayoría de la cartera de aplicaciones a la nube en oleadas, con más aplicaciones desplazadas a un ritmo más rápido en cada oleada. En esta fase, se utilizan las prácticas recomendadas y las lecciones aprendidas en las fases anteriores para implementar una fábrica de migración de equipos, herramientas y procesos con el fin de agilizar la migración de las cargas de trabajo mediante la automatización y la entrega ágil. Esta es la tercera fase de la [estrategia de migración de AWS](#).

fábrica de migración

Equipos multifuncionales que agilizan la migración de las cargas de trabajo mediante enfoques automatizados y ágiles. Los equipos de las fábricas de migración suelen incluir a analistas y propietarios de operaciones, empresas, ingenieros de migración, desarrolladores y DevOps profesionales que trabajan a pasos agigantados. Entre el 20 y el 50 por ciento de la cartera de aplicaciones empresariales se compone de patrones repetidos que pueden optimizarse mediante un enfoque de fábrica. Para obtener más información, consulte la [discusión sobre las fábricas de migración](#) y la [Guía de fábricas de migración a la nube](#) en este contenido.

metadatos de migración

Información sobre la aplicación y el servidor que se necesita para completar la migración. Cada patrón de migración requiere un conjunto diferente de metadatos de migración. Algunos ejemplos de metadatos de migración son la subred de destino, el grupo de seguridad y AWS la cuenta.

patrón de migración

Tarea de migración repetible que detalla la estrategia de migración, el destino de la migración y la aplicación o el servicio de migración utilizados. Ejemplo: realoje la migración a Amazon EC2 con AWS Application Migration Service.

Migration Portfolio Assessment (MPA)

Una herramienta en línea que proporciona información para validar el modelo de negocio para migrar a. Nube de AWS La MPA ofrece una evaluación detallada de la cartera (adecuación del tamaño de los servidores, precios, comparaciones del costo total de propiedad, análisis de los costos de migración), así como una planificación de la migración (análisis y recopilación de datos de aplicaciones, agrupación de aplicaciones, priorización de la migración y planificación de oleadas). La [herramienta MPA](#) (requiere iniciar sesión) está disponible de forma gratuita para todos los AWS consultores y consultores asociados de APN.

Evaluación de la preparación para la migración (MRA)

Proceso que consiste en obtener información sobre el estado de preparación de una organización para la nube, identificar sus puntos fuertes y débiles y elaborar un plan de acción para cerrar las brechas identificadas mediante el AWS CAF. Para obtener más información, consulte la [Guía de preparación para la migración](#). La MRA es la primera fase de la [estrategia de migración de AWS](#).

estrategia de migración

El enfoque utilizado para migrar una carga de trabajo a Nube de AWS Para obtener más información, consulte la entrada de las [7 R](#) de este glosario y consulte [Movilice a su organización para acelerar las migraciones a gran escala](#).

ML

[Consulte el aprendizaje automático.](#)

modernización

Transformar una aplicación obsoleta (antigua o monolítica) y su infraestructura en un sistema ágil, elástico y de alta disponibilidad en la nube para reducir los gastos, aumentar la eficiencia y aprovechar las innovaciones. Para obtener más información, consulte [Estrategia para modernizar las aplicaciones en el Nube de AWS](#).

evaluación de la preparación para la modernización

Evaluación que ayuda a determinar la preparación para la modernización de las aplicaciones de una organización; identifica los beneficios, los riesgos y las dependencias; y determina qué tan bien la organización puede soportar el estado futuro de esas aplicaciones. El resultado de la evaluación es un esquema de la arquitectura objetivo, una hoja de ruta que detalla las fases de desarrollo y los hitos del proceso de modernización y un plan de acción para abordar las brechas identificadas. Para obtener más información, consulte [Evaluación de la preparación para la modernización de las aplicaciones en el Nube de AWS](#).

aplicaciones monolíticas (monolitos)

Aplicaciones que se ejecutan como un único servicio con procesos estrechamente acoplados. Las aplicaciones monolíticas presentan varios inconvenientes. Si una característica de la aplicación experimenta un aumento en la demanda, se debe escalar toda la arquitectura. Agregar o mejorar las características de una aplicación monolítica también se vuelve más complejo a medida que crece la base de código. Para solucionar problemas con la aplicación, puede utilizar una arquitectura de microservicios. Para obtener más información, consulte [Descomposición de monolitos en microservicios](#).

MAPA

Consulte [la evaluación de la cartera de migración](#).

MQTT

Consulte [Message Queue Queue Telemetría](#) y Transporte.

clasificación multiclase

Un proceso que ayuda a generar predicciones para varias clases (predice uno de más de dos resultados). Por ejemplo, un modelo de ML podría preguntar “¿Este producto es un libro, un automóvil o un teléfono?” o “¿Qué categoría de productos es más interesante para este cliente?”.

infraestructura mutable

Un modelo que actualiza y modifica la infraestructura existente para las cargas de trabajo de producción. Para mejorar la coherencia, la fiabilidad y la previsibilidad, el AWS Well-Architected Framework recomienda el uso [de una infraestructura inmutable](#) como práctica recomendada.

O

OAC

[Consulte el control de acceso de origen.](#)

OAI

Consulte la [identidad de acceso de origen](#).

OCM

Consulte [gestión del cambio organizacional](#).

migración fuera de línea

Método de migración en el que la carga de trabajo de origen se elimina durante el proceso de migración. Este método implica un tiempo de inactividad prolongado y, por lo general, se utiliza para cargas de trabajo pequeñas y no críticas.

OI

Consulte [integración de operaciones](#).

OLA

Véase el [acuerdo a nivel operativo](#).

migración en línea

Método de migración en el que la carga de trabajo de origen se copia al sistema de destino sin que se desconecte. Las aplicaciones que están conectadas a la carga de trabajo pueden seguir

funcionando durante la migración. Este método implica un tiempo de inactividad nulo o mínimo y, por lo general, se utiliza para cargas de trabajo de producción críticas.

OPC-UA

Consulte [Open Process Communications: arquitectura unificada](#).

Comunicaciones de proceso abierto: arquitectura unificada (OPC-UA)

Un protocolo de comunicación machine-to-machine (M2M) para la automatización industrial. El OPC-UA proporciona un estándar de interoperabilidad con esquemas de cifrado, autenticación y autorización de datos.

acuerdo de nivel operativo (OLA)

Acuerdo que aclara lo que los grupos de TI operativos se comprometen a ofrecerse entre sí, para respaldar un acuerdo de nivel de servicio (SLA).

revisión de la preparación operativa (ORR)

Una lista de preguntas y las mejores prácticas asociadas que le ayudan a comprender, evaluar, prevenir o reducir el alcance de los incidentes y posibles fallos. Para obtener más información, consulte [Operational Readiness Reviews \(ORR\)](#) en AWS Well-Architected Framework.

tecnología operativa (OT)

Sistemas de hardware y software que funcionan con el entorno físico para controlar las operaciones, los equipos y la infraestructura industriales. En la industria manufacturera, la integración de los sistemas de TO y tecnología de la información (TI) es un enfoque clave para las transformaciones de [la industria 4.0](#).

integración de operaciones (OI)

Proceso de modernización de las operaciones en la nube, que implica la planificación de la preparación, la automatización y la integración. Para obtener más información, consulte la [Guía de integración de las operaciones](#).

registro de seguimiento organizativo

Un registro creado por el AWS CloudTrail que se registran todos los eventos para todos Cuentas de AWS los miembros de una organización AWS Organizations. Este registro de seguimiento se crea en cada Cuenta de AWS que forma parte de la organización y realiza un seguimiento de la actividad en cada cuenta. Para obtener más información, consulte [Crear un registro para una organización](#) en la CloudTrail documentación.

administración del cambio organizacional (OCM)

Marco para administrar las transformaciones empresariales importantes y disruptivas desde la perspectiva de las personas, la cultura y el liderazgo. La OCM ayuda a las empresas a prepararse para nuevos sistemas y estrategias y a realizar la transición a ellos, al acelerar la adopción de cambios, abordar los problemas de transición e impulsar cambios culturales y organizacionales. En la estrategia de AWS migración, este marco se denomina aceleración de personal, debido a la velocidad de cambio que requieren los proyectos de adopción de la nube. Para obtener más información, consulte la [Guía de OCM](#).

control de acceso de origen (OAC)

En CloudFront, una opción mejorada para restringir el acceso y proteger el contenido del Amazon Simple Storage Service (Amazon S3). El OAC admite todos los buckets de S3 Regiones de AWS, el cifrado del lado del servidor AWS KMS (SSE-KMS) y las solicitudes dinámicas PUT y DELETE dirigidas al bucket de S3.

identidad de acceso de origen (OAI)

En CloudFront, una opción para restringir el acceso y proteger el contenido de Amazon S3. Cuando utiliza OAI, CloudFront crea un principal con el que Amazon S3 puede autenticarse. Los directores autenticados solo pueden acceder al contenido de un bucket de S3 a través de una distribución específica. CloudFront Consulte también el [OAC](#), que proporciona un control de acceso más detallado y mejorado.

ORR

Consulte la revisión de [la preparación operativa](#).

OT

Consulte la [tecnología operativa](#).

VPC saliente (de salida)

En una arquitectura de AWS cuentas múltiples, una VPC que gestiona las conexiones de red que se inician desde una aplicación. La [arquitectura AWS de referencia de seguridad](#) recomienda configurar la cuenta de red con entradas, salidas e inspección VPCs para proteger la interfaz bidireccional entre la aplicación e Internet en general.

P

límite de permisos

Una política de administración de IAM que se adjunta a las entidades principales de IAM para establecer los permisos máximos que puede tener el usuario o el rol. Para obtener más información, consulte [Límites de permisos](#) en la documentación de IAM.

información de identificación personal (PII)

Información que, vista directamente o combinada con otros datos relacionados, puede utilizarse para deducir de manera razonable la identidad de una persona. Algunos ejemplos de información de identificación personal son los nombres, las direcciones y la información de contacto.

PII

Consulte la [información de identificación personal](#).

manual de estrategias

Conjunto de pasos predefinidos que capturan el trabajo asociado a las migraciones, como la entrega de las funciones de operaciones principales en la nube. Un manual puede adoptar la forma de scripts, manuales de procedimientos automatizados o resúmenes de los procesos o pasos necesarios para operar un entorno modernizado.

PLC

Consulte [controlador lógico programable](#).

PLM

Consulte la [gestión del ciclo de vida del producto](#).

policy

Un objeto que puede definir los permisos (consulte la [política basada en la identidad](#)), especifique las condiciones de acceso (consulte la [política basada en los recursos](#)) o defina los permisos máximos para todas las cuentas de una organización AWS Organizations (consulte la política de control de [servicios](#)).

persistencia políglota

Elegir de forma independiente la tecnología de almacenamiento de datos de un microservicio en función de los patrones de acceso a los datos y otros requisitos. Si sus microservicios tienen la misma tecnología de almacenamiento de datos, pueden enfrentarse a desafíos de

implementación o experimentar un rendimiento deficiente. Los microservicios se implementan más fácilmente y logran un mejor rendimiento y escalabilidad si utilizan el almacén de datos que mejor se adapte a sus necesidades. Para obtener más información, consulte [Habilitación de la persistencia de datos en los microservicios](#).

evaluación de cartera

Proceso de detección, análisis y priorización de la cartera de aplicaciones para planificar la migración. Para obtener más información, consulte la [Evaluación de la preparación para la migración](#).

predicate

Una condición de consulta que devuelve true o false, por lo general, se encuentra en una cláusula. WHERE

pulsar un predicado

Técnica de optimización de consultas de bases de datos que filtra los datos de la consulta antes de transferirlos. Esto reduce la cantidad de datos que se deben recuperar y procesar de la base de datos relacional y mejora el rendimiento de las consultas.

control preventivo

Un control de seguridad diseñado para evitar que ocurra un evento. Estos controles son la primera línea de defensa para evitar el acceso no autorizado o los cambios no deseados en la red. Para obtener más información, consulte [Controles preventivos](#) en Implementación de controles de seguridad en AWS.

entidad principal

Una entidad AWS que puede realizar acciones y acceder a los recursos. Esta entidad suele ser un usuario raíz para un Cuenta de AWS rol de IAM o un usuario. Para obtener más información, consulte Entidad principal en [Términos y conceptos de roles](#) en la documentación de IAM.

privacidad desde el diseño

Un enfoque de ingeniería de sistemas que tiene en cuenta la privacidad durante todo el proceso de desarrollo.

zonas alojadas privadas

Un contenedor que contiene información sobre cómo desea que Amazon Route 53 responda a las consultas de DNS de un dominio y sus subdominios dentro de uno o más VPCs. Para obtener más información, consulte [Uso de zonas alojadas privadas](#) en la documentación de Route 53.

control proactivo

Un [control de seguridad](#) diseñado para evitar el despliegue de recursos no conformes. Estos controles escanean los recursos antes de aprovisionarlos. Si el recurso no cumple con el control, significa que no está aprovisionado. Para obtener más información, consulte la [guía de referencia de controles](#) en la AWS Control Tower documentación y consulte [Controles proactivos](#) en Implementación de controles de seguridad en AWS.

gestión del ciclo de vida del producto (PLM)

La gestión de los datos y los procesos de un producto a lo largo de todo su ciclo de vida, desde el diseño, el desarrollo y el lanzamiento, pasando por el crecimiento y la madurez, hasta el rechazo y la retirada.

entorno de producción

Consulte [el entorno](#).

controlador lógico programable (PLC)

En la fabricación, una computadora adaptable y altamente confiable que monitorea las máquinas y automatiza los procesos de fabricación.

encadenamiento rápido

Utilizar la salida de una solicitud de [LLM](#) como entrada para la siguiente solicitud para generar mejores respuestas. Esta técnica se utiliza para dividir una tarea compleja en subtareas o para refinar o ampliar de forma iterativa una respuesta preliminar. Ayuda a mejorar la precisión y la relevancia de las respuestas de un modelo y permite obtener resultados más detallados y personalizados.

seudonimización

El proceso de reemplazar los identificadores personales de un conjunto de datos por valores de marcadores de posición. Laseudonimización puede ayudar a proteger la privacidad personal. Los datosseudonimizados siguen considerándose datos personales.

publish/subscribe (pub/sub)

Un patrón que permite las comunicaciones asíncronas entre microservicios para mejorar la escalabilidad y la capacidad de respuesta. Por ejemplo, en un [MES](#) basado en microservicios, un microservicio puede publicar mensajes de eventos en un canal al que se puedan suscribir otros microservicios. El sistema puede añadir nuevos microservicios sin cambiar el servicio de publicación.

Q

plan de consulta

Serie de pasos, como instrucciones, que se utilizan para acceder a los datos de un sistema de base de datos relacional SQL.

regresión del plan de consulta

El optimizador de servicios de la base de datos elige un plan menos óptimo que antes de un cambio determinado en el entorno de la base de datos. Los cambios en estadísticas, restricciones, configuración del entorno, enlaces de parámetros de consultas y actualizaciones del motor de base de datos PostgreSQL pueden provocar una regresión del plan.

R

Matriz RACI

Véase [responsable, responsable, consultado, informado \(RACI\)](#).

RAG

Consulte [Retrieval Augmented Generation](#).

ransomware

Software malicioso que se ha diseñado para bloquear el acceso a un sistema informático o a los datos hasta que se efectúe un pago.

Matriz RASCI

Véase [responsable, responsable, consultado, informado \(RACI\)](#).

RCAC

Consulte control de [acceso por filas y columnas](#).

réplica de lectura

Una copia de una base de datos que se utiliza con fines de solo lectura. Puede enrutar las consultas a la réplica de lectura para reducir la carga en la base de datos principal.

rediseñar

Ver [7 Rs](#).

objetivo de punto de recuperación (RPO)

La cantidad de tiempo máximo aceptable desde el último punto de recuperación de datos. Esto determina qué se considera una pérdida de datos aceptable entre el último punto de recuperación y la interrupción del servicio.

objetivo de tiempo de recuperación (RTO)

La demora máxima aceptable entre la interrupción del servicio y el restablecimiento del servicio.

refactorizar

Ver [7 Rs.](#)

Región

Una colección de AWS recursos en un área geográfica. Cada una Región de AWS está aislada e independiente de los demás para proporcionar tolerancia a las fallas, estabilidad y resiliencia. Para obtener más información, consulte [Regiones de AWS Especificar qué cuenta puede usar](#).

regresión

Una técnica de ML que predice un valor numérico. Por ejemplo, para resolver el problema de “¿A qué precio se venderá esta casa?”, un modelo de ML podría utilizar un modelo de regresión lineal para predecir el precio de venta de una vivienda en función de datos conocidos sobre ella (por ejemplo, los metros cuadrados).

volver a alojar

Consulte [7 Rs.](#)

versión

En un proceso de implementación, el acto de promover cambios en un entorno de producción.

trasladarse

Ver [7 Rs.](#)

redefinir la plataforma

Ver [7 Rs.](#)

recompra

Ver [7 Rs.](#)

resiliencia

La capacidad de una aplicación para resistir las interrupciones o recuperarse de ellas. [La alta disponibilidad](#) y la [recuperación ante desastres](#) son consideraciones comunes a la hora de planificar la resiliencia en el. Nube de AWS Para obtener más información, consulte [Nube de AWS Resiliencia](#).

política basada en recursos

Una política asociada a un recurso, como un bucket de Amazon S3, un punto de conexión o una clave de cifrado. Este tipo de política especifica a qué entidades principales se les permite el acceso, las acciones compatibles y cualquier otra condición que deba cumplirse.

matriz responsable, confiable, consultada e informada (RACI)

Una matriz que define las funciones y responsabilidades de todas las partes involucradas en las actividades de migración y las operaciones de la nube. El nombre de la matriz se deriva de los tipos de responsabilidad definidos en la matriz: responsable (R), contable (A), consultado (C) e informado (I). El tipo de soporte (S) es opcional. Si incluye el soporte, la matriz se denomina matriz RASCI y, si la excluye, se denomina matriz RACI.

control receptivo

Un control de seguridad que se ha diseñado para corregir los eventos adversos o las desviaciones con respecto a su base de seguridad. Para obtener más información, consulte [Controles receptivos](#) en Implementación de controles de seguridad en AWS.

retain

Consulte [7 Rs](#).

jubilarse

Ver [7 Rs](#).

Generación aumentada de recuperación (RAG)

Tecnología de [inteligencia artificial generativa](#) en la que un máster [hace referencia](#) a una fuente de datos autorizada que se encuentra fuera de sus fuentes de datos de formación antes de generar una respuesta. Por ejemplo, un modelo RAG podría realizar una búsqueda semántica en la base de conocimientos o en los datos personalizados de una organización. Para obtener más información, consulte [Qué es](#) el RAG.

rotación

Proceso de actualizar periódicamente un [secreto](#) para dificultar el acceso de un atacante a las credenciales.

control de acceso por filas y columnas (RCAC)

El uso de expresiones SQL básicas y flexibles que tienen reglas de acceso definidas. El RCAC consta de permisos de fila y máscaras de columnas.

RPO

Consulte el [objetivo del punto de recuperación](#).

RTO

Consulte el [objetivo de tiempo de recuperación](#).

manual de procedimientos

Conjunto de procedimientos manuales o automatizados necesarios para realizar una tarea específica. Por lo general, se diseñan para agilizar las operaciones o los procedimientos repetitivos con altas tasas de error.

S

SAML 2.0

Un estándar abierto que utilizan muchos proveedores de identidad (IdPs). Esta función permite el inicio de sesión único (SSO) federado, de modo que los usuarios pueden iniciar sesión AWS Management Console o llamar a las operaciones de la AWS API sin tener que crear un usuario en IAM para todos los miembros de la organización. Para obtener más información sobre la federación basada en SAML 2.0, consulte [Acerca de la federación basada en SAML 2.0](#) en la documentación de IAM.

SCADA

Consulte el [control de supervisión y la adquisición de datos](#).

SCP

Consulte la [política de control de servicios](#).

secreta

Información confidencial o restringida, como una contraseña o credenciales de usuario, que almacene de forma cifrada. AWS Secrets Manager Se compone del valor secreto y sus metadatos. El valor secreto puede ser binario, una sola cadena o varias cadenas. Para obtener más información, consulta [¿Qué hay en un secreto de Secrets Manager?](#) en la documentación de Secrets Manager.

seguridad desde el diseño

Un enfoque de ingeniería de sistemas que tiene en cuenta la seguridad durante todo el proceso de desarrollo.

control de seguridad

Barrera de protección técnica o administrativa que impide, detecta o reduce la capacidad de un agente de amenazas para aprovechar una vulnerabilidad de seguridad. Existen cuatro tipos principales de controles de seguridad: [preventivos](#), [de detección](#), con [capacidad](#) de [respuesta](#) y [proactivos](#).

refuerzo de la seguridad

Proceso de reducir la superficie expuesta a ataques para hacerla más resistente a los ataques. Esto puede incluir acciones, como la eliminación de los recursos que ya no se necesitan, la implementación de prácticas recomendadas de seguridad consistente en conceder privilegios mínimos o la desactivación de características innecesarias en los archivos de configuración.

sistema de información sobre seguridad y administración de eventos (SIEM)

Herramientas y servicios que combinan sistemas de administración de información sobre seguridad (SIM) y de administración de eventos de seguridad (SEM). Un sistema de SIEM recopila, monitorea y analiza los datos de servidores, redes, dispositivos y otras fuentes para detectar amenazas y brechas de seguridad y generar alertas.

automatización de la respuesta de seguridad

Una acción predefinida y programada que está diseñada para responder automáticamente a un evento de seguridad o remediarlo. Estas automatizaciones sirven como controles de seguridad [detectables](#) o [adaptables](#) que le ayudan a implementar las mejores prácticas AWS de seguridad. Algunos ejemplos de acciones de respuesta automatizadas incluyen la modificación de un grupo de seguridad de VPC, la aplicación de parches a una EC2 instancia de Amazon o la rotación de credenciales.

cifrado del servidor

Cifrado de los datos en su destino, por parte de quien Servicio de AWS los recibe.

política de control de servicio (SCP)

Política que proporciona un control centralizado de los permisos de todas las cuentas de una organización en AWS Organizations. SCPs defina barreras o establezca límites a las acciones que un administrador puede delegar en usuarios o roles. Puede utilizarlas SCPs como listas de permitidos o rechazados para especificar qué servicios o acciones están permitidos o prohibidos. Para obtener más información, consulte [las políticas de control de servicios](#) en la AWS Organizations documentación.

punto de enlace de servicio

La URL del punto de entrada de un Servicio de AWS. Para conectarse mediante programación a un servicio de destino, puede utilizar un punto de conexión. Para obtener más información, consulte [Puntos de conexión de Servicio de AWS](#) en Referencia general de AWS.

acuerdo de nivel de servicio (SLA)

Acuerdo que aclara lo que un equipo de TI se compromete a ofrecer a los clientes, como el tiempo de actividad y el rendimiento del servicio.

indicador de nivel de servicio (SLI)

Medición de un aspecto del rendimiento de un servicio, como la tasa de errores, la disponibilidad o el rendimiento.

objetivo de nivel de servicio (SLO)

[Una métrica objetivo que representa el estado de un servicio, medido mediante un indicador de nivel de servicio.](#)

modelo de responsabilidad compartida

Un modelo que describe la responsabilidad que compartes con respecto a la seguridad y AWS el cumplimiento de la nube. AWS es responsable de la seguridad de la nube, mientras que usted es responsable de la seguridad en la nube. Para obtener más información, consulte el [Modelo de responsabilidad compartida](#).

SIEM

Consulte [la información de seguridad y el sistema de gestión de eventos](#).

punto único de fallo (SPOF)

Una falla en un único componente crítico de una aplicación que puede interrumpir el sistema.

SLA

Consulte el acuerdo [de nivel de servicio](#).

SLI

Consulte el indicador de [nivel de servicio](#).

SLO

Consulte el objetivo de nivel de [servicio](#).

split-and-seed modelo

Un patrón para escalar y acelerar los proyectos de modernización. A medida que se definen las nuevas funciones y los lanzamientos de los productos, el equipo principal se divide para crear nuevos equipos de productos. Esto ayuda a ampliar las capacidades y los servicios de su organización, mejora la productividad de los desarrolladores y apoya la innovación rápida. Para obtener más información, consulte [Enfoque gradual para modernizar las aplicaciones en el. Nube de AWS](#)

SPOF

Consulte el [punto único de falla](#).

esquema en forma de estrella

Estructura organizativa de una base de datos que utiliza una tabla de hechos grande para almacenar datos medidos o transaccionales y una o más tablas dimensionales más pequeñas para almacenar los atributos de los datos. Esta estructura está diseñada para usarse en un [almacén de datos](#) o con fines de inteligencia empresarial.

patrón de higo estrangulador

Un enfoque para modernizar los sistemas monolíticos mediante la reescritura y el reemplazo gradual de las funciones del sistema hasta que se pueda dismantelar el sistema heredado. Este patrón utiliza la analogía de una higuera que crece hasta convertirse en un árbol estable y, finalmente, se apodera y reemplaza a su host. El patrón fue [presentado por Martin Fowler](#) como una forma de gestionar el riesgo al reescribir sistemas monolíticos. Para ver un ejemplo con la aplicación de este patrón, consulte [Modernización gradual de los servicios web antiguos de Microsoft ASP.NET \(ASMX\) mediante contenedores y Amazon API Gateway](#).

subred

Un intervalo de direcciones IP en la VPC. Una subred debe residir en una sola zona de disponibilidad.

supervisión, control y adquisición de datos (SCADA)

En la industria manufacturera, un sistema que utiliza hardware y software para monitorear los activos físicos y las operaciones de producción.

cifrado simétrico

Un algoritmo de cifrado que utiliza la misma clave para cifrar y descifrar los datos.

pruebas sintéticas

Probar un sistema de manera que simule las interacciones de los usuarios para detectar posibles problemas o monitorear el rendimiento. Puede usar [Amazon CloudWatch Synthetics](#) para crear estas pruebas.

indicador del sistema

Una técnica para proporcionar contexto, instrucciones o pautas a un [LLM](#) para dirigir su comportamiento. Las indicaciones del sistema ayudan a establecer el contexto y las reglas para las interacciones con los usuarios.

T

etiquetas

Pares clave-valor que actúan como metadatos para organizar los recursos. AWS Las etiquetas pueden ayudarle a administrar, identificar, organizar, buscar y filtrar recursos. Para obtener más información, consulte [Etiquetado de los recursos de AWS](#).

variable de destino

El valor que intenta predecir en el ML supervisado. Esto también se conoce como variable de resultado. Por ejemplo, en un entorno de fabricación, la variable objetivo podría ser un defecto del producto.

lista de tareas

Herramienta que se utiliza para hacer un seguimiento del progreso mediante un manual de procedimientos. La lista de tareas contiene una descripción general del manual de

procedimientos y una lista de las tareas generales que deben completarse. Para cada tarea general, se incluye la cantidad estimada de tiempo necesario, el propietario y el progreso.

entorno de prueba

[Consulte entorno.](#)

entrenamiento

Proporcionar datos de los que pueda aprender su modelo de ML. Los datos de entrenamiento deben contener la respuesta correcta. El algoritmo de aprendizaje encuentra patrones en los datos de entrenamiento que asignan los atributos de los datos de entrada al destino (la respuesta que desea predecir). Genera un modelo de ML que captura estos patrones. Luego, el modelo de ML se puede utilizar para obtener predicciones sobre datos nuevos para los que no se conoce el destino.

puerta de enlace de tránsito

Un centro de tránsito de red que puede usar para interconectar sus VPCs redes con las locales. Para obtener más información, consulte [Qué es una pasarela de tránsito](#) en la AWS Transit Gateway documentación.

flujo de trabajo basado en enlaces troncales

Un enfoque en el que los desarrolladores crean y prueban características de forma local en una rama de característica y, a continuación, combinan esos cambios en la rama principal. Luego, la rama principal se adapta a los entornos de desarrollo, preproducción y producción, de forma secuencial.

acceso de confianza

Otorgar permisos a un servicio que especifique para realizar tareas en su organización AWS Organizations y en sus cuentas en su nombre. El servicio de confianza crea un rol vinculado al servicio en cada cuenta, cuando ese rol es necesario, para realizar las tareas de administración por usted. Para obtener más información, consulte [AWS Organizations Utilización con otros AWS servicios](#) en la AWS Organizations documentación.

ajuste

Cambiar aspectos de su proceso de formación a fin de mejorar la precisión del modelo de ML. Por ejemplo, puede entrenar el modelo de ML al generar un conjunto de etiquetas, incorporar etiquetas y, luego, repetir estos pasos varias veces con diferentes ajustes para optimizar el modelo.

equipo de dos pizzas

Un DevOps equipo pequeño al que puedes alimentar con dos pizzas. Un equipo formado por dos integrantes garantiza la mejor oportunidad posible de colaboración en el desarrollo de software.

U

incertidumbre

Un concepto que hace referencia a información imprecisa, incompleta o desconocida que puede socavar la fiabilidad de los modelos predictivos de ML. Hay dos tipos de incertidumbre: la incertidumbre epistémica se debe a datos limitados e incompletos, mientras que la incertidumbre aleatoria se debe al ruido y la aleatoriedad inherentes a los datos. Para más información, consulte la guía [Cuantificación de la incertidumbre en los sistemas de aprendizaje profundo](#).

tareas indiferenciadas

También conocido como tareas arduas, es el trabajo que es necesario para crear y operar una aplicación, pero que no proporciona un valor directo al usuario final ni proporciona una ventaja competitiva. Algunos ejemplos de tareas indiferenciadas son la adquisición, el mantenimiento y la planificación de la capacidad.

entornos superiores

Ver [entorno](#).

V

succión

Una operación de mantenimiento de bases de datos que implica limpiar después de las actualizaciones incrementales para recuperar espacio de almacenamiento y mejorar el rendimiento.

control de versión

Procesos y herramientas que realizan un seguimiento de los cambios, como los cambios en el código fuente de un repositorio.

Emparejamiento de VPC

Una conexión entre dos VPCs que le permite enrutar el tráfico mediante direcciones IP privadas. Para obtener más información, consulte [¿Qué es una interconexión de VPC?](#) en la documentación de Amazon VPC.

vulnerabilidad

Defecto de software o hardware que pone en peligro la seguridad del sistema.

W

caché caliente

Un búfer caché que contiene datos actuales y relevantes a los que se accede con frecuencia. La instancia de base de datos puede leer desde la caché del búfer, lo que es más rápido que leer desde la memoria principal o el disco.

datos templados

Datos a los que el acceso es infrecuente. Al consultar este tipo de datos, normalmente se aceptan consultas moderadamente lentas.

función de ventana

Función SQL que realiza un cálculo en un grupo de filas que se relacionan de alguna manera con el registro actual. Las funciones de ventana son útiles para procesar tareas, como calcular una media móvil o acceder al valor de las filas en función de la posición relativa de la fila actual.

carga de trabajo

Conjunto de recursos y código que ofrece valor comercial, como una aplicación orientada al cliente o un proceso de backend.

flujo de trabajo

Grupos funcionales de un proyecto de migración que son responsables de un conjunto específico de tareas. Cada flujo de trabajo es independiente, pero respalda a los demás flujos de trabajo del proyecto. Por ejemplo, el flujo de trabajo de la cartera es responsable de priorizar las aplicaciones, planificar las oleadas y recopilar los metadatos de migración. El flujo de trabajo de la cartera entrega estos recursos al flujo de trabajo de migración, que luego migra los servidores y las aplicaciones.

GUSANO

Mira, [escribe una vez, lee muchas](#).

WQF

Consulte el [marco AWS de calificación de la carga](#) de trabajo.

escribe una vez, lee muchas (WORM)

Un modelo de almacenamiento que escribe los datos una sola vez y evita que los datos se eliminen o modifiquen. Los usuarios autorizados pueden leer los datos tantas veces como sea necesario, pero no pueden cambiarlos. Esta infraestructura de almacenamiento de datos se considera [inmutable](#).

Z

ataque de día cero

Un ataque, normalmente de malware, que aprovecha una vulnerabilidad de [día cero](#).

vulnerabilidad de día cero

Un defecto o una vulnerabilidad sin mitigación en un sistema de producción. Los agentes de amenazas pueden usar este tipo de vulnerabilidad para atacar el sistema. Los desarrolladores suelen darse cuenta de la vulnerabilidad a raíz del ataque.

aviso de tiro cero

Proporcionar a un [LLM](#) instrucciones para realizar una tarea, pero sin ejemplos (imágenes) que puedan ayudar a guiarla. El LLM debe utilizar sus conocimientos previamente entrenados para realizar la tarea. La eficacia de las indicaciones cero depende de la complejidad de la tarea y de la calidad de las indicaciones. [Consulte también las indicaciones de pocos pasos](#).

aplicación zombi

Aplicación que utiliza un promedio de CPU y memoria menor al 5 por ciento. En un proyecto de migración, es habitual retirar estas aplicaciones.

Las traducciones son generadas a través de traducción automática. En caso de conflicto entre la traducción y la version original de inglés, prevalecerá la version en inglés.