

Bases de Datos y Bases de Datos Avanzadas

modelamiento de Bases de datos.

→ **entidad**: cualquier sujeto medible y contable.
Contienen atributos.

→ **relacion**: asociación entre entidades

→ **cardinalidad**: nivel de asociación de la relación

- 1-1 ; 1-n ; n-1 ; n-m.

→ **dimension de medida**:

- bit a byte.

- 1 byte = 8 bits

- 1 Mbyte = 1024 kbytes

- 1 Kbyte = 1024 bytes

- 1 Gbyte = 1024 Mbytes

- 1 Tbyte = 1024 Gbytes

→ **llaves**:

atributos que identifican una única ocurrencia de una entidad o relación

- **super llaves** → **llaves candidatas**; **llaves primarias** (identifican las tuplas en una tabla de forma única)

- **llaves foráneas** → **dependen de otra tabla**

→ clave que tiene valores que coinciden con los valores de clave primaria de otra relación

→ **Tipos de datos**:

tipo dato	tamaño almacenamiento	rango
• int	4 bytes	-2.147.483.648 - 2.147.483.647
• VARCHAR	n + 1 bytes	0-255 caracteres
• year	1 byte	1901 - 2185 (año)
• date	3 bytes	1/01/1001 - 31/12/9999
• DATETIME	8 bytes	1/01/1001 a las 0:0:0 - 31/12/9999 a las 23:59:59
• TEXT	longitud + 2 bytes	0-255 caracteres
• TIMESTAMP	4 bytes	1/01/1970 - 2037

→ **Operadores unarios**

• **seleccion** → (llamar a una tabla)

• **proyeccion** → (seleccionar atributos)

• **renombramiento** → (renombrar atributos y entidades)

→ **operadores binarios**

- Resta, producto, asociación, division, union, interseccion

→ **estructura de una consulta JOIN**

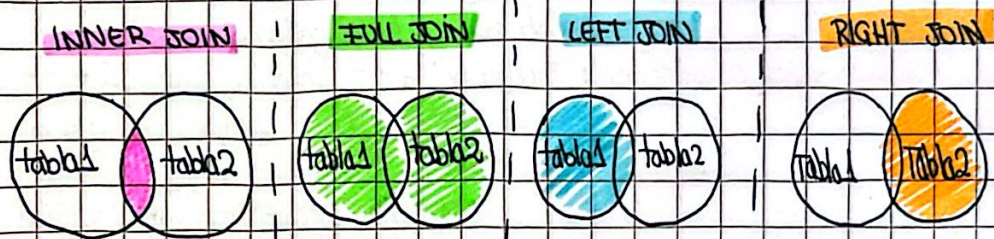
SELECT <atributos>

FROM <tabla> (**INNER** | **LEFT** | **RIGHT** | **FULL**) **JOIN**

<tabla2> **on** <condition>

→ estructura de una consulta

SELECT <atributo>
FROM <tabla>
WHERE <condition>



TRAER RESULTADOS DE TABLAS QUE CUMPLIRAN CON LA COMPARACION ENTRE COLUMNAS

INFORMACION RESULTADOS COMPARACION DE TABLAS QUE CUMPLAN CON LAS COMPARACION ENTRE COLUMNAS Y RESULTADOS DE TABLAS (DERECHA, IZQUIERDA)

IGUAL QUE LEFT JOIN PERO CON LA TABLA DERECHA.

→ FUNCIONES INTEGRADAS

- MAX (<atributo>): retorna el valor mayor.
- MIN (<atributo>): retorna el valor menor.
- SUM (<atributo>): suma de valores (entre ellos).
- AVG (<atributo>): promedio entre valores.
- ORDER BY (<atributo>): ordena los resultados de una consulta según el atributo indicado (lo hace de manera ascendente / DESC → descendente).
- COUNT (<atributo>): contador.
- GROUP BY (<atributo>): agrupa las filas que tienen los mismos valores.
- LIMIT <número>: obtener el número exacto de resultados.
- BETWEEN (<valor1> AND <valor2>): rango inclusivo para un atributo.

SELECT * FROM <tabla> WHERE ...

- NOT BETWEEN (<valor1> AND <valor2>): negativo del BETWEEN.
- IN / WHERE <atributo> IN (<valor1>, <valor2>): permite buscar argumentos específicos de una lista.
- NOT IN: negativo del IN.

BETWEEN / NOT BETWEEN

SELECT *
FROM <TABLA>
WHERE <atributo> BETWEEN / NOT BETWEEN
<valor1> AND <valor2>

IN / NOT IN

SELECT <atributo>
FROM <tabla>
WHERE <atributo> IN / NOT IN
(<valor1>, <valor2>)

NORMALIZACIÓN

• conjunto de reglas obtenidas luego el paso del modelo entidad-relación las cuales busca:

- Evitar la redundancia de datos
- Evitar problemas de actualización de los datos en las tablas.
- Proteger la integridad de los datos.
- Agilizar las consultas o las actualizaciones o las inserciones o eliminaciones.

1FN o Primera Forma Normal

→ Todas las tablas tienen atributos conformados por un solo valor.

Ej:

nombre	rut	masкота
Hola 1	11111-1	bruno, cachupin
Hola 2	2222-2	coliflor, malehin



nombre	rut	masкота
hola 1	11111-1	bruno
hola 1	11111-1	cachupin
hola 2	2222-2	coliflor
hola 2	2222-2	malehin

2FN o Segunda Forma Normal

→ Cumple el 1FN

→ Todos los atributos (que no sean parte de la clave candidata) entregan información de la clave completa.

Ej:

id mascota	nombre	rut	masкота	raza
113	hola 1	11111-1	bruno	labrador
114	hola 1	11111-1	cachupin	Puddle
225	hola 2	2222-2	coliflor	bulldog
227	hola 2	2222-2	malehin	fox terrier



id mascota	masкота	raza
113	bruno	labrador
114	cachupin	puddle
225	coliflor	bulldog
227	malehin	fox terrier

id	rut
113	11111-1
114	11111-1
225	2222-2
227	2222-2

nombre	rut
hola 1	11111-1
hola 2	2222-2

PROARTE®

3FN o Tercera Forma Normal

→ Cumple con 2FN

→ Todos los atributos dependen directamente de la llave primaria

Ej:

idmascota	mascota	raza	idmascota	rut	nombre	rut
113	bruno	labrador	113	1111-1	Hola1	1111-1
114	coliflor	fox terrier	114	1111-1	Hola2	2222-2
225	malehin	bulldogg	225	2222-2		

idmascota	mascota
113	bruno
114	coliflor
225	malehin

mascota	raza
bruno	labrador
coliflor	fox terrier
malehin	bulldogg

nombre	rut
hola1	1111-1
hola2	2222-2

idmascota	rut
113	1111-1
114	1111-1
225	2222-2

Bases de datos relacionales

- Maneja información estructurada
- Almacena la información en tablas (filas y columnas).
- Establece relaciones entre tablas usando llaves.
- Utiliza SQL como el lenguaje de administración y manejo de la bdd.
- Orientado a sistemas transaccionales.
- Garantiza las propiedades ACID (aislamiento, consistencia, ~~o~~ durabilidad).

ventajas:

- asegura la consistencia de los datos
- organiza los datos, elimina anomalías (duplicidad de datos, inconsistencias)
- dispone una gran variedad de herramientas que simplifican el manejo de la bdd.

desventajas:

- orientada a trabajar en una sola máquina.
- escalabilidad limitada.
- esquema rígido de datos, difícil de modificar.
- rendimiento decrece en la medida que la bdd crece.

Bases de datos no relacionales

- maneja información no estructurada
- almacena info usando estructuras de datos.
- los datos no siguen un patrón específico.
- almacena cualquier tipo de dato.
- orientada al manejo de grandes volúmenes de datos.
- datos no siguen un patrón específico.
- sacrifica consistencia en pro de la eficiencia.
- orientada a sistemas de gestión (OLAP).

Tipos:

- Key-values stores: diccionario no estructurado.
- Document Database: almacena documentos (JSON, BSON, XML)
- Graph Database: almacena datos en nodos.

ventajas:

- almacenamiento flexible de los datos
- permite manejar grandes volúmenes de datos.
- posibilita la escalabilidad horizontal.
- simplifica la modificación del esquema de datos

desventajas:

- No existe lenguaje estándar de administrar la base de datos
- No asegura la consistencia de los datos.
- Requiere una gran cantidad de almacenamiento debido a la redundancia de datos.

Vistas:

- forma diferente de visualizar los datos físicos, se materializa mediante JOINs entre tablas.
- visión lógica distinta de los datos físicos reales.
- garantiza la independencia de datos y aporta confidencialidad, y sigue funcionando con normalidad si se hacen consultas.
- se gestiona lo que en los caracteres empleamos para definirla (solo).

ventajas:

- aporta confidencialidad a la BDD.
- si se estructura la BDD, alguna normalización, se puede mantener la estructura de las consultas en la BDD.
- se crean mediante el comando view.
- Los vistas (lógica) tiene una representación real (física) en una tabla.

VISTAS MATERIALIZADAS:

- Concentración de resultado de la consulta que genera la vista, es decir, la vista se crea y los resultados de esta vista quedan guardados en el disco, en la BDD.
- se actualiza sola.
- la vista materializada no se actualiza, se queda con los datos con la que se creó, es decir, no se puede borrar, ni agregar, etc.

PLAN DE EJECUCIÓN

- se genera mediante el comando explain.

componentes:

- 1º seq scan: secuencia de búsqueda por bloques.
 - primer cost: costo del sistema para iniciar el proceso.
 - costo: valor arbitrario.
 - segundo cost: costo del sistema para terminar el proceso.
 - rows: cantidad de filas que se recorren.
 - width: cuantos bytes contiene cada fila.
- 2º Group Key = tabla de hash.
- 3º Nested loop = ciclo interno.
- 4º Index Only Scan = búsqueda logarítmica por que se utiliza un index.
 - solo búsqueda por índice con el Index Only Scan.
 - por cada una de las filas, se encuentra una, se demora un cost, encuentra solo 1 fila, por ende row = 1 y width = 8.
- 5º Hash Aggregate = (cost = (cuanto cuesta), (total)).

Tipos de búsqueda:

- **secuencial**: búsqueda por bloque.
- **bitmap index**: sistema guarda en memoria el recorrido necesario de lo que se pide sin ejecutarlos.
- **Index only**: recomiendo el árbol, no necesito ir a buscar los datos.
- **Index**: búsqueda por índice, navegar el B-tree y buscar los datos que necesito.
- **Bitmap heap**: Sistema guarda los datos en un heap.

Instrucciones de la BDD como funcionan:

SELECT

- Recibe un conjunto de condiciones que pueden ser de selección o de asociación. Luego parte por las de selección con el fin de optimizar la consulta.
- Busca las filas que cumplan con la selección.
- Asocia esas filas con las otras tablas.
- Esto implica una cantidad "N" de accesos (bloques secuencial como bloque del árbol del índice).
- Obtiene el resultado solicitado.

INSERT

Caso 1: sin índice

- Trae a memoria el bloque último de la tabla, si el bloque tiene espacio, agrega la fila a ese bloque.
- Si el bloque no tiene espacio, genera un nuevo bloque y agrega la fila.
- Graba el bloque en el dispositivo de almacenamiento final de la tabla. +1 acceso.

Caso 2: con índice

- Se busca el nodo en donde se debe agregar el valor del atributo.
- Una vez ubicado el nodo se trae a memoria, si hay espacio en ese nodo se agrega el nuevo valor.
- Si no hay espacio, se realiza un block split (nodo se divide en 2, cada ~~parte~~ parte con la mitad de los valores por lo tanto se inserta el valor más la dirección a memoria).

DELETE

- busca filas que cumplen con la condición de selección.
- marcar en el bloque la eliminación de la fila y luego grabar esto en el dispositivo de almacenamiento.
- Si hay índice, se busca el nodo en donde se debe eliminar el valor y la dirección del atributo; elimina este valor y graba el nodo nuevamente en el dispositivo de almacenamiento.

UPDATE

- Buscar filas que cumplan con la condición de selección.
- Llevar a memoria el bloque que contiene la fila.
- La fila en cuestión se marca como eliminada y se crea otra fila con la actualización.
- Trae el último bloque en la tabla en cuestión, si hay espacio se agrega a la fila a ese bloque; si no hay espacio se genera un nuevo bloque para insertar la fila actualizada.
- Si existe un índice, se busca el nodo en cuestión, en el nodo se actualiza la dirección, el espacio será el mismo, así que no hay problema, luego se graba en el dispositivo de almacenamiento.

RAIDS: redundant array of independent disks

- conjuntos de discos duros que se instalan en un rack, de manera que permite a la bdd ampliar su capacidad.
- mejora el rendimiento de la base de datos.

TIPS DE RAID

• RAID 0: Data Striping, striped volume

2 discos	
A1	A2
A3	A4
A5	A6
A7	A8

- volumen dividido, capacidad suma de discos.
- no tolera fallos, se guarda alternadamente.
- duplica la ~~capacidad~~ capacidad de almacenamiento, para lectura se tienen dos bloques o recuperar reduciendo el tiempo de acceso a la mitad.

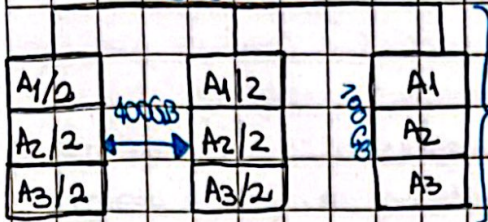
• RAID 1:

2 discos	
A1	A1
A2	A2
A3	A3

- disco espejo, no se pierde información y no se aumenta el espacio.
- se tienen datos redundantes.
- protección de datos, ya que se tienen los mismos datos.
- en lectura se pueden leer ambos discos, recuperando y reduciendo el tiempo.

RAID 3: Disco extra de paridad "data striping".

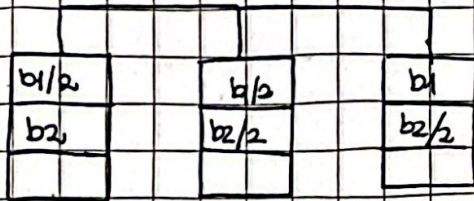
3 discos



- disco extra de paridad dedicado
- guarda la paridad de cada bloque
- si se cae uno de los 3 discos por medio del disco de paridad y del 2 restantes se recupera el perdido.

- capacidad: 2 de los discos - discos de paridad.
- soporta el fallo de un disco.
- duplica la capacidad de almacenamiento, es mas lento y recupera un disco en caso de perderse uno mediante el disco de chequeo y el restante.
- es mas lento que RAID 0 y RAID 4.

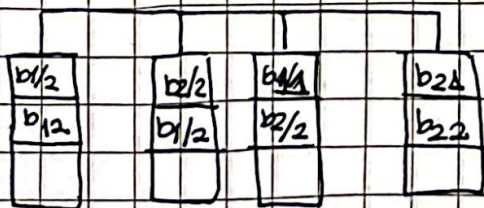
RAID 5: distribuido con paridad



3 discos

- paridad distribuida en todos los discos.
- capacidad suma de discos - disco de paridad.
- mejora lectura, penaliza escritura.
- soporta el fallo de un disco.
- funciona = que raid 3 pero el bloque de chequeo se puede ir moviendo en los discos.

RAID 6:



4 discos

- funciona de la misma manera que RAID 5 pero el bloque de chequeo se va moviendo y son dos bloques de chequeo.
- se puede recuperar un disco, si se llegan a perder dos.

RAID 10:



4 discos

- se tiene un RAID 1 repetitivo, o sea un espejo.
- se lee simultaneamente de 4 discos, por ende, el tiempo se divide en 4 o sea menor.
- menos discos, menos espacio desperdiciado.
- se pierde menos espacio dependiendo el RAID.

TABLESPACE

- espacio en el disco en donde se almacena todos los objetos que crea la base de datos.
- para el sistema operativo es solo un archivo que ocupa un determinado espacio extensible para la base de datos "estodo" y postgres puede ver la estructura de este archivo.
- limitado por el tamaño del disco.
- no se debe llenar uno para crear otro, el sistema necesita espacio para procesos internos.
- se crea mediante `tablespace` (comando).