

DISK, FILES

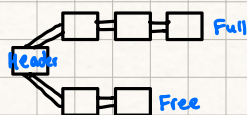
Read: Disk $\xrightarrow{\text{Page}}$ RAM
 Write: RAM $\xrightarrow{\text{Page}}$ Disk
 Files SSD (Flash) $\hookrightarrow$ Platter
 Pages $\hookrightarrow$ Cells $\hookrightarrow$ Sequential
 Records $\hookrightarrow$ Fast random reads
 $\hookrightarrow$ Failure after erasure

FILES

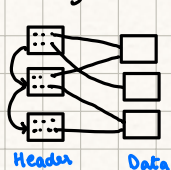
Unordered Heap Files

$\hookrightarrow$ No order

$\hookrightarrow$ 1. Linked List



$\hookrightarrow$ 2. Page Directory



• Entries in header page:
 < Pointer, Free Space >
 $\hookrightarrow$
 • Insertion faster

$\downarrow$ Searching: full scan (N I/Os)

Sorted Files

$\hookrightarrow$ Pages ordered

$\hookrightarrow$ Recorded sorted by Key(s)

$\uparrow$ Searching: $\log N$ I/O

$\downarrow$ Insertion: $\log N + N$ I/Os
 (req. shuffling if Keys unsorted,
 read all data pages)

SQL

Tables make up relational databases

NULL: falsey, any type

GROUPING

$\hookrightarrow$ summarize cols [SUM, AVG, MAX, COUNT]

$\hookrightarrow$ ignores NULL except COUNT(*)

check select statement w grouping

$\hookrightarrow$ where $\rightarrow$ grouping $\rightarrow$ having

$\hookrightarrow$ if any grouping done, only
 SELECT grouped/agg cols

$\hookrightarrow$ default ASC

$\hookrightarrow$ UNION removes duplicates

SELECT [DISTINCT] <columns> ⑤

FROM <table> ①

[INNER/LEFT/RIGHT/FULL] JOIN <table> FROM

ON <predicate>

[WHERE <predicate>] ②

[GROUP BY <columns>] ③

[HAVING <predicate>] ④

[ORDER BY <column list>] ⑥

[LIMIT <# rows>]; ⑦

LIKE %: 0+ chars, _ any char

Subqueries:

WHERE

WITH

UNION- distinct in either

INTERSECT - in both

ALL - all subquery

ANY - meet condition

ANY - any one true

JOINS

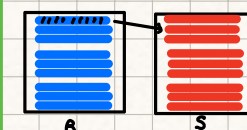
[R][S] $\rightarrow$ # pages

[R][S] $\rightarrow$ # records/page

exclude final
 writes to
 disk

SNLJ

Every record $\rightarrow$ Every Record

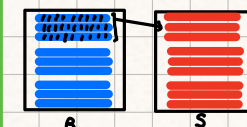


for r_i in R:
 for s_j in S:
 if $\theta(r_i, s_j)$
 yield

$$[R] + [R][S]$$

Do both ones, take the min one

PNLJ



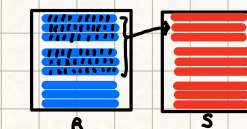
for each P_R in R:
 for each P_S in S:
 for each row r in P_R :
 for each row s in P_S :
 <...>

$$[R] + [R][S]$$

Keep smaller relation as outer loop

SNLJ, PNLJ use 3 buffers (2 input + 1 output)

BNLJ



for each block of (B-2) pages C_R :
 for each P_S in S:
 for each row r in C_R :
 for each row s in P_S :
 <...>

$$[R] + \left\lceil \frac{[R]}{B-2} \right\rceil [S]$$

(B-2) buffers for smaller one, 1 for other, 1 output

INLJ

for each record r_i in R:

for each record s_j in S where $\theta(s_i, r_j)$ is true:
 <yield>

$$[R] + |R| \times (\text{cost to lookup in S})$$

SORT-MERGE JOIN:

1. Sort both R and S

2. Merge matching tuples

```
while (True) {
  while (r < s) { advance r }
  while (r > s) { advance s }
  mark
  while (r == s) {
    while (r == s) {
      yield <r, s>
      advance
    }
    reset s to mark
    advance r
  }
}
```

Don't need to materialize
 sorted relation

In final merge pass of
 sorting both, stream
 to SMT

$\rightarrow$ SAVES $2([R] + [S])$

Need:
 # runs in
 last merge
 of R + # runs in
 last merge
 of S
 $\leq (B-1)$

Optimised: Total - $2([S] + [R])$

RELATIONAL ALGEBRA

Projection π [select] takes columns

Selection σ [where] filters rows

Union $\cup$ combine diff relations, remove duplicates } MUST have
 Set Diff $-$ rows in table 1 not in table 2 } same attributes

Group By γ [GROUP BY]

Intersection $\cap$ rows in both tables

Cross Product $\times$ 1 tuple for every possible pair

Alias ρ [AS] renames attributes, for join

Join $\bowtie$ default natural join

$\wedge$ (AND) $\vee$ (OR)

can't have duplicates (tuples)

After π , other cols DON'T exist

$$S_1 - (S_1 - S_2) = S_1 \cap S_2 = S_1 \bowtie S_2$$

B+ TREES

- Order d , always balanced
- Must have sorted $d \leq x \leq 2d$ entries
- Inner Nodes have max $(2d+1)$ pointers
- Only leaf nodes have records

Max values: $2df(2d+1)^h$

INSERTION

- Find leaf node, add $\langle K, R \rangle$ in order
- Overflow ($L > 2d$)
 - $L_1(d)$ $L_2(d+1)$
 - If L leaf, copy to parent
 - If L inner, PUSH to parent
 - Adjust pointers
 - Recurse

DELETION

- Just delete from leaf. NEVER from inner

STORAGE

- Alt 1 (By Val): Leaf contains data $\langle K, R \rangle$ $\downarrow$ can't support multiple indices, sorted
- Alt 2 (By Ref): Contains pointers $\langle K, \# \text{Page} \# \text{Record} \rangle$ $\uparrow$ multiple indices
- Alt 3 (By List of Ref): List of pointers $\langle K, [...] \rangle$ $\uparrow$ multiple records, same leaf node entry

CLUSTERING

Unclustered: $\sim 1/0$ per record

Clustered: $\sim 1/0$ per page

COUNTING I/Os

- Read root $\rightarrow$ leaf
- Read data page (un vs clustered)
- Write data page
- Write index page

BULK LOADING

- Sort data
- Fill leaf to f
- Add pointer to parent
- Parent Overflow:
 - $L_1(d)$ $L_2(d+1)$
 - Move L_2 first entry to parent
- Adjust pointers

SELECT: best case: doesn't exist, so only root $\rightarrow$ leaf

Check if pages overflow

can be built on any type of column

full scan = # data pages

Index Search: $h+1 = \log_F(\# \text{leaves}) + 1$

Bulkloading makes use of temporal locality

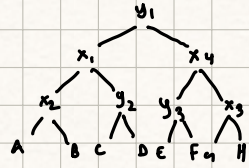
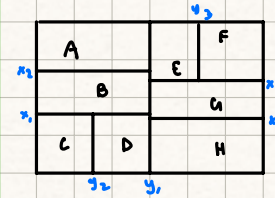
SPATIAL INDICES

Brute force search = # calc = # points

K-d tree search

K-d tree that organizes points in a K-dim space

Balanced binary tree that splits on dimensions



EQUALITY SEARCH: $\log(BR) + 2$ $B: \# \text{Blocks}$ $R: \# \text{Records / block}$

RANGE SEARCH: $\log(BR) + 1 + 2 \cdot \# \text{pages}$

DELETION: $\log BR + 4$ INSERTION: $\log BR + 4$ $\rightarrow$ search, read leaf, write data, write index

Searching:

$\rightarrow$ Use BST to find block its in

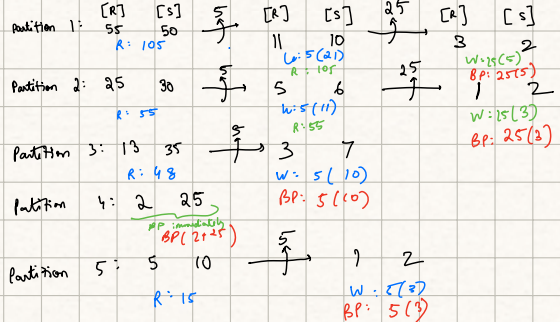
$\rightarrow$ Backtrack if other neighbour block is closer

JOINS PT 2

GRACE HASH JOIN:

- Partition Data
 - $\rightarrow$ Partition R, S into $(B-1)$ partitions
 - $\rightarrow$ Recurse until either $\leq (B-2)$
- Build and Probe
 - $\rightarrow$ Build in memory hash table for smaller one

Pass 1: 500



COSTS

	Heap	Sorted	Clustered
Scan All	B	B	$\frac{1}{f} \cdot B$
Equality	$B/2$	$\log_2 B$	$\log_F \left(\frac{BR}{E} \right) + 2$
Range	B	$\log_2 B + \text{pages}$	$\log_F \left(\frac{BR}{E} \right) + 3 \cdot \# \text{pages}$
Insert	2B	$\log_2 B + B$	$\log_F \left(\frac{BR}{E} \right) + 4$
Delete	$\frac{B}{2} + 1$	$\log_2 B + B$	$\log_F \left(\frac{BR}{E} \right) + 4$

B: # Data Blocks
R: # Records / Block

F: Fanout
E: # Entries per leaf

} $\propto D$
time per scan

BUFFER MANAGEMENT

- Buffer frame holds 1 data page perfectly
- Metadata Table: Frame ID, Page ID, Dirty Bit, Pin Count
 - memory addr
 - pg on frame
 - modified
 - # requests
- LRU: evict oldest used and that is UNPINNED
- CLOCK: LRU approx, uses ref bit $\uparrow$ time
 - Skip pinned until unpinned + ref bit 0 found
 - If ref=0, evict, write, read new page + set ref bit to 1
 - Move clock hand to next frame
 - Hint: set ref bit to 1, DON'T move hand
- MRU: evict most recently used $\uparrow$ sequential scans

Requester sets dirty page

Page could have multiple pins

MRU > LRU sequential, = if $P < F$

SORTING

B buffer pages, $(B-1)$ input, 1 output

I/Os: $2N \cdot \left(1 + \log_{B-1} \frac{N}{B} \right)$

Stop when 1 run

$N=1000, B=11$

Pass 0: $1000/11 = 91$ runs, 11 each

Pass 1: $91/10 = 10$ runs, 110 each

Pass 2: $10/10 = 1$ run, 1100 each

pages to sort in p passes: $B'(B-1)^{p-1} \geq N$

I/Os does NOT depend on # B pages

HASHING

Hash into $(B-1)$ partitions, recurse until $\leq B$, 1 for streaming in

Group By + Duplicate Values

Hash func need to be distinct + 1 for conquer phase

Data skew $\downarrow$

Hash w/o recursive part: $B(B-1)$

B buffers, p passes $\Rightarrow (B-1)^{p-1} \cdot B \geq \# \text{max pages}$

10 buffers, 3 passes $\Rightarrow (10-1)^{3-1} \cdot 10 = 810$

part conquer

all unique values will be in 1 partition ALWAYS

QUERY OPTIMIZATION

Streaming operator: Work to produce each tuple

Blocking operator: Consume entire input THEN output generated

Selectivity Estimation

Pred	Selectivity
$C = V$	$\frac{1}{ C }, \frac{1}{ V }$
$C_1 = C_2$	$\frac{1}{\max(C_1, C_2)}$
$C \leq V$ INT	$\frac{V - \min_c}{\max_c - \min_c + 1} + \frac{1}{ C }$
$C \geq V$ INT	$\frac{\max_c - V}{\max_c - \min_c + 1} + \frac{1}{ C }$
$C \leq V$ FLOAT	$\frac{V - \min_c}{\max_c - \min_c}$
$C \geq V$ FLOAT	$\frac{\max_c - V}{\max_c - \min_c}$
$P_1 \text{ AND } P_2$	$S(P_1) \cdot S(P_2)$
$P_1 \text{ OR } P_2$	$S(P_1) + S(P_2) - S(P_1)S(P_2)$
NOT P	$1 - S(P)$

$B = 5$ $[P] = 50$, $[T] = 100$
 $P_u = 30$ $T_u = 40$
mat both: $(50 \times 30) + (100 \times 40) + 30 + \lceil \frac{30}{3} \rceil 40$
no mat: $50 + \lceil \frac{30}{3} \rceil 100$

Atomicity: all or none

Consistency: start & end consistent

Isolation: isolated from other Xact

Durability: commit should persist

Serial Equivalent:

Same Xact, order, state

View Equivalent: Same initial read, dependent read, winning writes

2PL: Ensures conflict serializable

1. S before X 2. No acquiring after releasing

Strict 2PL: Prevents cascading aborts

1. Release all locks together

Lock Management: Hash table {Resource: (Granted, Mode, Queue)}

Deadlock Avoidance (T_i waiting for T_j)

1. Wait Die: T_i higher priority $\rightarrow T_i$ waits
 T_j lower priority $\rightarrow T_j$ aborts

2. Wound Wait: T_i higher priority $\rightarrow T_j$ aborts
 T_j lower priority $\rightarrow T_i$ waits

Lock Granularity

Compatibility							Parent						
Y/N	NL	IS	IX	S	SIX	X	Y/N	NL	IS	IX	S	SIX	X
NL	✓	✓	✓	✓	✓	✓	NL	✓	✓	✓	✓	✓	✓
IS	✓	✓	✓	✓	✓	✗	IS	✓	✓	✓	✓	✓	✗
IX	✓	✓	✓	✗	✗	✗	IX	✓	✓	✓	✓	✓	✗
S	✓	✓	✗	✓	✗	✗	S	✓	✗	✗	✗	✗	✗
SIX	✓	✓	✗	✗	✓	✗	SIX	✓	✗	✗	✗	✓	✗
X	✓	✗	✗	✗	✗	✓	X	✓	✗	✗	✗	✗	✓

RIES

1. Analysis: Recreate DPT, Xact Table

Start at begin checkpoint

Not End: Add to Xact, last LSN = LSN

Update/CLR Undo: Add to DPT, recLSN = LSN

Commit/Abort: Change status in Xact

End: Remove from Xact

At end, any committing txn: put END record, remove from Xact

running txn: put ABORT record, change to aborting

Durability

2. REDO: Reconstruct state

Start at smallest recLSN

REDO ALL Update/CLR.

DON'T if

1. Page not in DPT

2. recLSN > LSN

3. pageLSN > LSN

DISTRIBUTED TRANSACTIONS

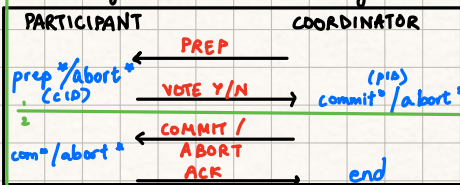
Distributed Locking

↳ Every node has own lock table (independent)

↳ UNION waits-for graph for distributed deadlock

2PC

↳ 2PC guarantees serializability



↳ Log + flush THEN send

↳ Presumed Abort: DON'T flush abort

↳ Commit ONLY happens when co logs commit

2PC Recovery

1. Participant Recovering

↳ No prepare: abort

↳ Prepare: ask coord

↳ Abort I: abort + vote NO

↳ Abort II: abort + send ack

↳ Commit: send ack

2. Coordinator Recovering

↳ No commit: abort

↳ End: Nothing

↳ Abort: rerun phase 2

PARALLEL QUERY PROCESSING

Architectures

1. Shared memory: CPUs share mem + disk

2. Shared disk: CPUs share disk

3. Shared nothing

Partitioning schemes

↳ Sharding: 1 data page on 1 machine

↳ Replication: 1 data page on many machines

1. Range Partitioning ↑ Key lookup, range

Each machine gets range of values

2. Hash Partitioning ↑ Key lookup

Each machine gets hashed values

3. Round Robin ↑ parallelisation ↓ all req activation

Each machine gets same #pages

Parallelism

↳ Inter-query ↑ throughput

Each machine dif queries

↳ Intra-query

One query across machines

↳ Intra-operator

One operator fast as possible

↳ Inter-operator

Operators in parallel

1. Pipeline (pass to parent)

2. Bushy Tree (dif branches TT)

Network Cost

Data sent over network $\frac{pk(n-1)}{n}$

Parallel Sorting: Range partition + local sort

#passes = $1 + (1 + \log_{B-1} \frac{N}{MB})$

Parallel Hashing: Hash partition + local hash

Parallel SMJ: Range partition BOTH + local SMJ

#passes = $1 + 1 + \#passes_R + \#passes_S + 1 + 1$

I/Os = $[R] + [S] + (1 + \log_{B-1} \frac{[R]}{B}) + (1 + \log_{B-1} \frac{[S]}{B}) + [R] + [S]$

Parallel GHJ: Hash partition BOTH + local GHJ

Symmetric Hash Join:

1. Build 2 hash tables

2. R/S record arrives, probe S/R

Hierarchical Aggregation:

↳ COUNT → m: count c: sum

↳ AVG → m: sum + count

↳ c: add + divide

1. Steal/No-steal (Atomicity)

Can modified, uncommitted Xact be flushed?

↳ Steal: Yes (UNDO)

↳ No-steal: No (bad performance)

2. Force/No-force (Durability)

Are modified pages forced to disk on commit?

↳ Force: Yes

↳ No-Force: No (REDO)

3. UNDO:

Start at end, move up

UNDOes all updates for

Running/Aborting Xacts

↳ Write CLR UNDO Tx: LSN prevLSN, UNDO NEXT LSN

↳ Once all undone, write end record

↳ Remove Xact

Undo Logging: Steal/Force

Dirty pages to disk B4 commit log

Flush pages B4 commit

<T, A, Prev>

If T com/ab → completed else write to disk

Redo Logging: NF/NS

commit B4 flush: redo if crash after commit but not flush

<T, A, next val> ignore uncom txn com: write new v

NOSQL

Online Transaction Processing (OLTP):

- high no. of transactions ex large no. of users
- workload: frontend (social net, online stores)
- queries: simple lookups, updates

Online Analytical Processing (OLAP):

- read only workloads
- large amount of data
- large no. of joins + aggregations

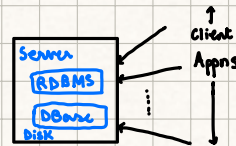
	OLTP	OLAP
main read main write used	small no. of rec/ query rand. access, low latency end user/customer via web app	aggregate over large no. of rec bulk import or event stream internal analysis decision support
data	current state	history
size	GB - TB	TB - PB

Architectures:

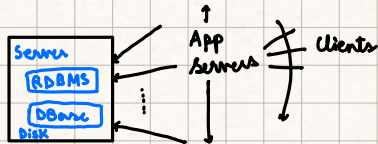
1. One-Tier Architecture



2. Two-Tier Architecture



3. Three-Tier Architecture



Partitioning/Sharding: Efficient for write-heavy workloads

Replication: Scale, resilient, extensive down time

CAP Theorem:

- Consistency: 2 clients simultaneous get same view
- Availability: Every request → Not a faulty response
- Partition Tolerance: Continue to operate despite drop/delay

Impossible for $> 2/3$ to be satisfied

↑ Consistency ⇒ Error ↑ Availability ⇒ Stale