

ADBS EX 5

appears to user like centralized DBMS

Distributed Database nodes over network, possible no homogeneity
 +: no single point of failure, isolate faults, replace nodes, load balancing

shared Memory: common memory - memory contention
shared Disk: common disk, SAM, - congestion

shared Nothing: linear speed up, linear scale up
Fragmentation: relation $\rightarrow$ smaller fragments **Allocation**: distrib. frag, replication

Replication: several copies, repl. factor: # of copies
Distribution Transparency: location, fragmentation, migration, concurrency, failure

Failures in DBMS: server failure, link f., message f., Network partition - difficult
Availability: Probability System is operational **Reliability**: Sys. is failure free

Fault tolerance: Failures happen, detect failures, no inconsistent data, cont. operation
Goals of Fragmentation/Replication: Data Locality, $\downarrow$ communication cost

small sets $\rightarrow$ improved efficiency, load balancing, - find compromise
Fragmentation: Horizontal/Sharding: $R = R_1 \cup R_2 \cup R_3$ / Vertical/Columns $R = R_1 \bowtie R_2 \bowtie R_3$

Non Redundancy/Disjointness $R_1 \cap R_2 = \emptyset$ / $PK(R) = attr(R_1) \cap attr(R_2) = \emptyset$

Granularity of Frag.: not too coarse/fine, workload aware, local frag. criterion
 + **Replication**: availability, read performance - consistency, concurrency

Primary - Secondary Repl.: write to primary + no concurrency - bottleneck
Multi Primary: synchronize servers + write availability - consistency, concurrency

Atomicity Consistency Isolation Durability
Distributed Concurrency Control deadlock detection, **Distr. Recovery** after fail.

Distinguished Copy Lock, **Voting** majority/timeout - more traffic, complex
 at Primary Site for all copies, Primary copy on different Servers

Dist. Deadlock Detection: **Centralized**: send waits-for-graph to central site
Hierarchical: send info to parent **Time-out**: abort after timeout

Two Phase Commit: coordinator force write log
 prepare subordinate yes/no write comm. Log
 coordinator force write log
 subordinate yes/no write comm. Log
 coordinator write end log

Join: $\pi_B(R) \rightarrow S, S' = \pi_B(R) \bowtie S \rightarrow R, R \bowtie S' = R \bowtie S \rightarrow$ to site 3

+ of **RDBMS**: formal foundation (logic, rel. algebra), Standard Query Language
 Concurrency, multi user, integration Database

- of **RDBMS**: schema known in advance, structured data, column type, split into Tables, rows
Want: nested Structure/JSON, graphs, schema evolution, sparse data, many nulls

$\rightarrow$ application DBS $\rightarrow$ REST API - flexible, scale out
NoSQL: distributed, schemaless, non-relational, no joins, horizontally scalable

Basically **Available Soft State Eventual Consistency** - no transactions
Aggregates: Key Value opaque, Document Store JSON, Column Store, Graph Relations, Nodes

Arbitrary Content Structure combination of C. combination
Consistency - like R/W on single copy **Availability** **Partition Tolerance** - continue operation

AP or CP, CA is impractical, **PA/EC** Latency Consistency, **PA/EC** or **PC/EL**
ROWA read one - write all: $W = N, R = 1$, **Majority**: $W \geq \lfloor N/2 \rfloor + 1, R \geq \lfloor N/2 \rfloor$

Lamport Clock $G_j = \max(C_j, C_i) + 1$ 0 - 1 - 2 - 3 - 4 - 5
Weak Clock: $e_1 \rightarrow e_2$ then $C(e_1) < C(e_2)$ 0 - 1 - 2 - 3 - 4 - 5 - 6

Strong Clock: $e_1 \rightarrow e_2 \leftrightarrow C(e_1) < C(e_2)$ 0 - 1 - 2 - 3 - 4 - 5 - 6
Vector Clock: $VC_j[k] = \max(VC_j[k], VC_i[k])$, then own counter $VC_j[j] = VC_j[j] + 1$

Key Value: value is array of bytes - meaning only known to app, scalable, available
Session: User $\rightarrow$ Session Data, **Logs**: Date $\rightarrow$ Logfile, **Ad**: Campaign ID $\rightarrow$ Ad content

Riak: Erlang, masterless, available/fault tolerant, tag+link values in Metadata
 When Bucket/namespace is created: Set N # of Replicas, R-read Quorum, W-write Quorum

Document Store: understands value XML, JSON, BSON no joins $\rightarrow$ denormalize
MongoDB for log data, product catalogue, CMS data, collection ~ table, document ~ row

1:n rel: one - millions $\rightarrow$ foreign key one - thousands $\rightarrow$ array of references one - few $\rightarrow$ embed
 n:m rel: three collections like SQL, or 2 with reference array embedded

Denormalization: - updates not atomic, + avoid join when high read to write ratio

MongoDB primary-secondary writes on primary → propagate to secondary, error → vote new primary
write concern - when is write successful, majority = $\lfloor N/2 \rfloor + 1$, **read concern** e.g. majority
Sharding by Shard Key - Hashed Key or Ranged Key, access through Router

Graph DB Edges Directed, optimized for Graph Traversal, Nodes + Edges have proper
Use Case: Customer-products, companies-ownership, people-places, social network
+ edges stored as pointers - no slow down if big - no index, no schema, easy to adapt/change

Graph in RDBMS: Node * Table, Edge Table, Node property T., edge prop. T.
or one Table <id, propertyKey, value> eg <4, Name, 'Bob'>, <5, since, 1990>

Transform ERD to Graph: TABLE → Node Label, Row → Node, column → N. prop., FK → Relations

Neo4J: ACID compliant, several primaries, secondaries replicate asynchronously
locking, eventual consistency, sharding difficult - different secondaries cache in RAM

Cypher: declarative, pattern matching, () node, (p) match var p against node
(:Person) node Label 'Person', (p:Person {name: 'Tom Hanks'})
→ /<-- directed rel. --> undirected rel., -[:ACTS_IN]-> rel. type ACTS_IN
Match (a)--(c) where (a)-[works_in]->(c:City) Return a.name
-[:ACTED_IN*1..3]- 1 to 3 hops Match p=(foo)--(bar) return nodes(p)/rels(p)
p = Shortest Path (x)--(y) Return nodes(p), length(p), relationships(p)
Creates erstellt Merge updates/creates
Match ... Set p.age = 23 Match ... Delete p

Aggregate: avg, sum, max, min, count(*), collect(p.name) - List of values, order by
Constraints: uniqueness, existence, property key

Column stores: vertical fragments / no rows, virtual id or explicit id SIMD
+ aggregation, only needed columns - less IO, CPU heavy - less disk, vectorized
Explicit ID: larger Data → more IO **Virtual ID**: lookup with offset: start + n * width

Compression Run length encoding: <value, # repetitions> can be worse than no comp.

Bit Vector: n vectors for n distinct values, 1 if = value, else 0

Dictionary Encoding: Dictionary with values, column has dict keys
Null suppression for sparse data, vector has 0 for Null, separate dictionary
++ if I can operate on compressed data, e.g. RLE sum = value * # repetitions

Distinct values still good: RLE 20, Bit Vec. 30, Dictionary 60-100, LZ 100-200

Date Column → RLE, **Month** → Bit Vector, **Address** → Dictionary

Late Materialization: operate on individual columns, join late

Updates: expensive on compressed columns - decompress / update / compress
→ write to write store, flush periodically to read store, read/merge from both

Wide Column: ^{HBase, Cassandra} Google Bigtable: scale, continuous updates, lots of IO

Distributed, multi-level, sparse map: Row key - column - timestamp

ordered by key, grouped into Tablets, columns in column family

column contents c.family anchor: cnsi.com qualifier ↑ timestamp
row com.cnn.rwn <html>..Lp7.. cnn.com

Lookup: return most recent values / return timestamp range