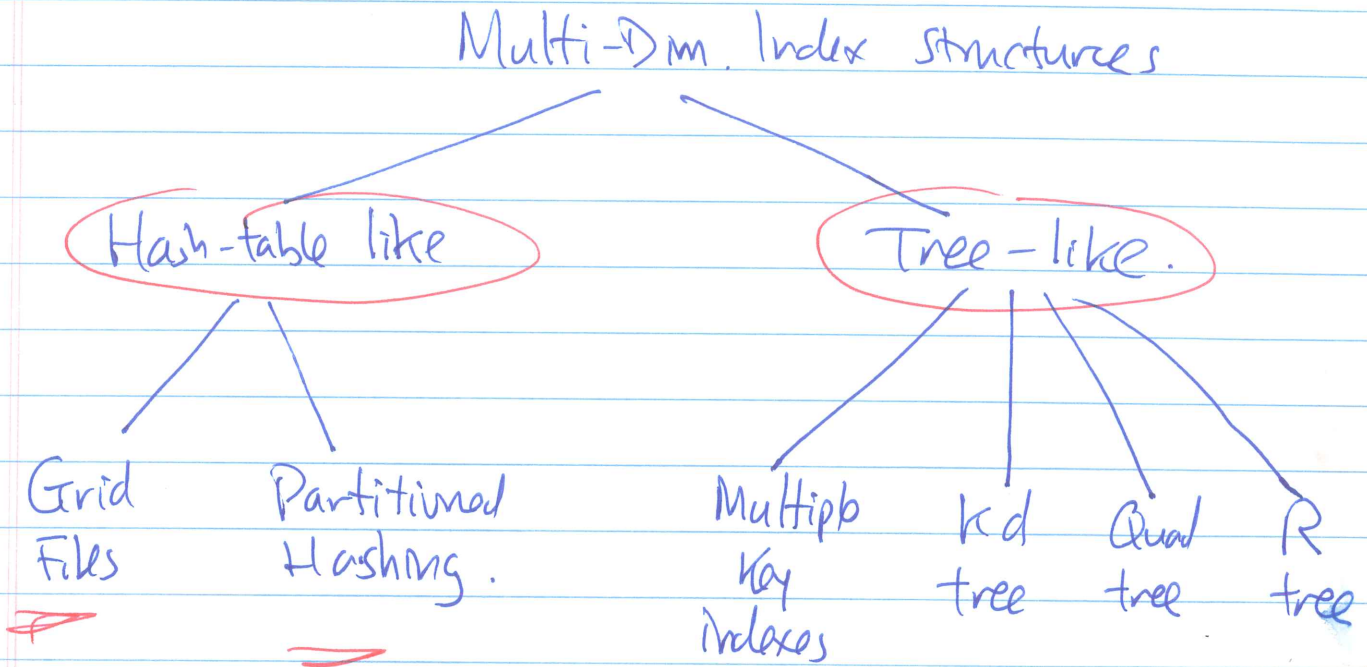


Grid Index Files + Partitioned Hash Functions

6

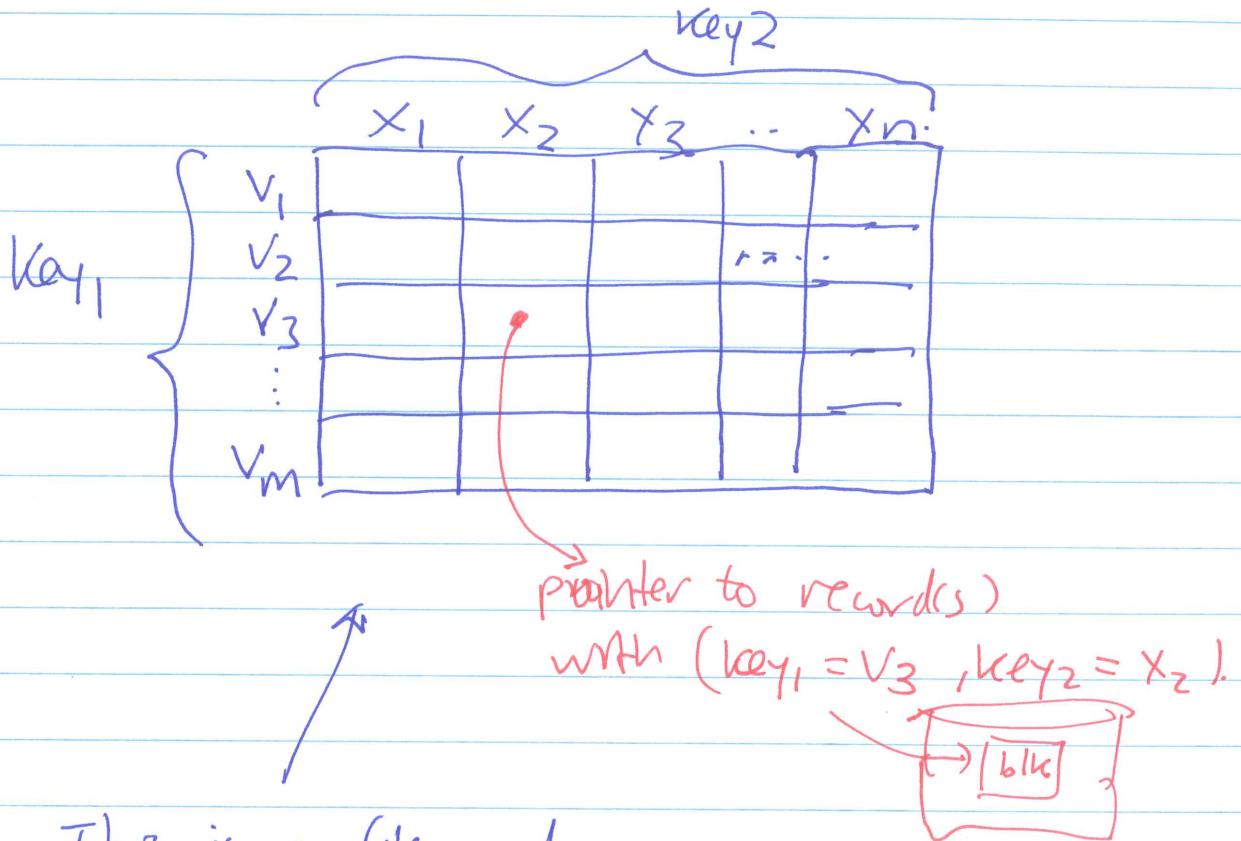
Multi-Dim. Index Structures:



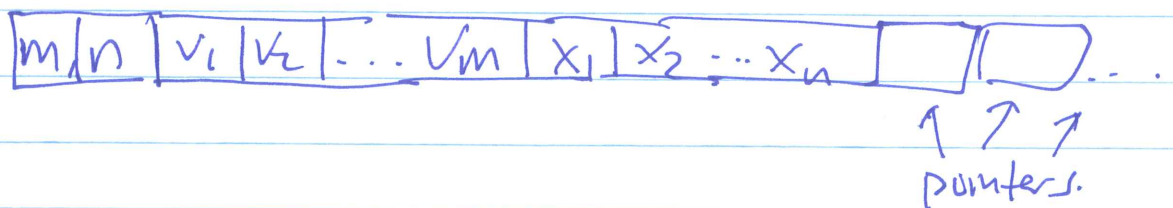
↑ ↑
This set of notes

Grid Index (Grid Files)

- A Grid Index (File) is ^{an index} organized into a 2-dimensional structure:

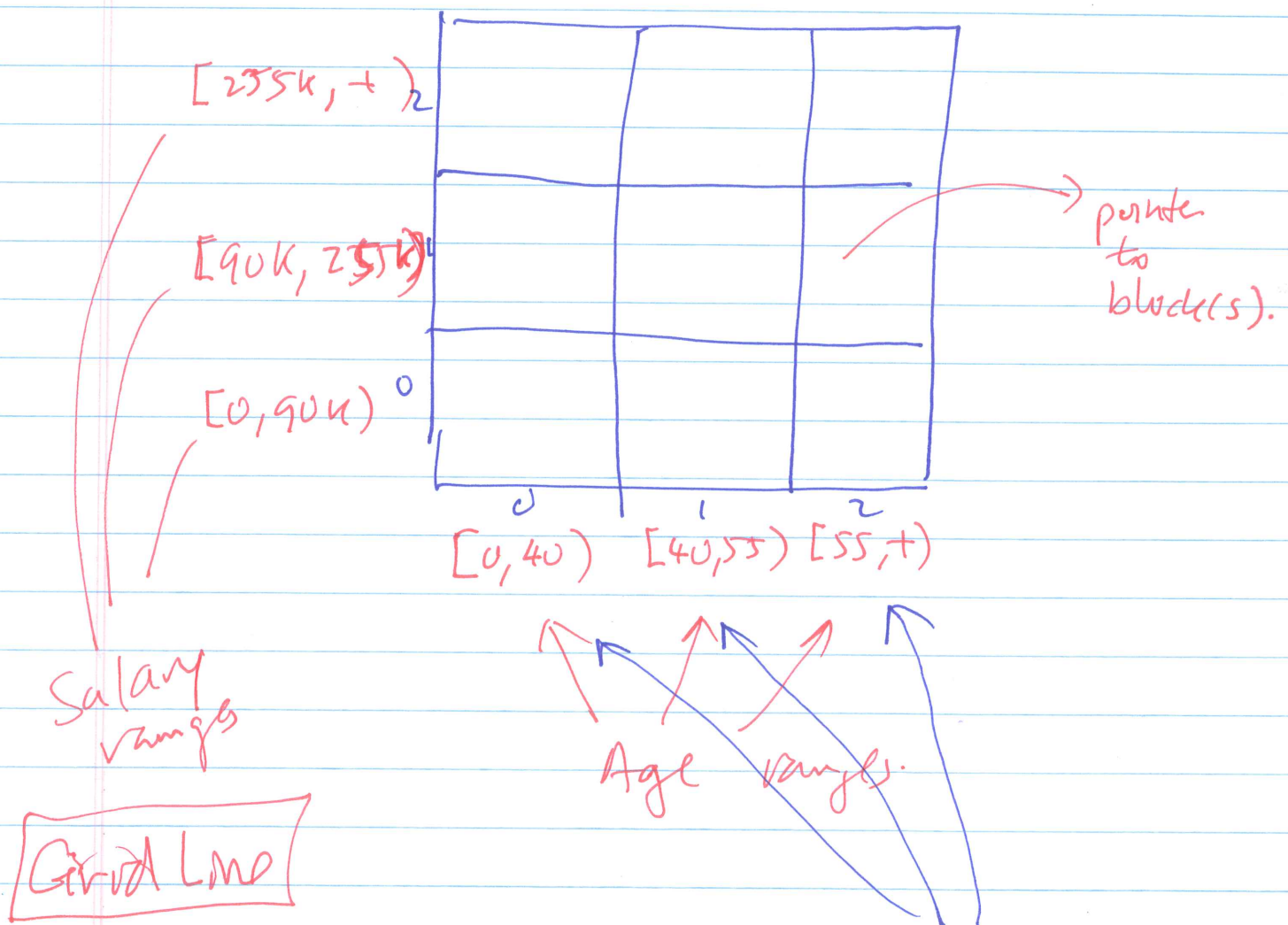


This is a file and stored eg. as follows:



In general, the keys are
ranges
or "buckets".

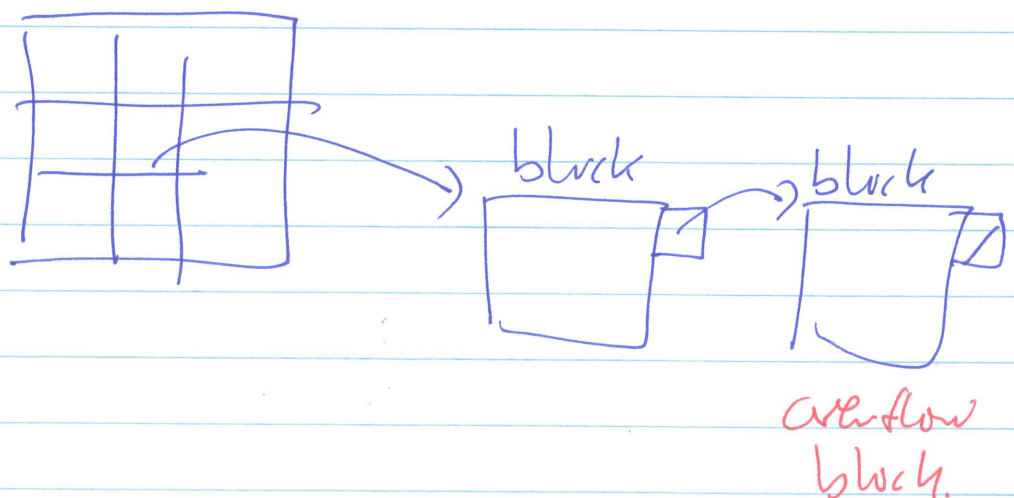
E.g.:



- Ranges are called "Grid Lines"

Grid Index file.

- Uses an array whose dimension is the same as for the data file.
- Store - for each dimension - the list of values at which the "grid lines" occur.
- Uses overflow blocks: blocks.



Example index:

- Suppose we have a bunch of records of people buying jewelry
- Assume the only relevant attributes are:

(~~name~~, address, age, salary)

↑ ↑
important properties
to categorize buyers

- We want to index the information on:

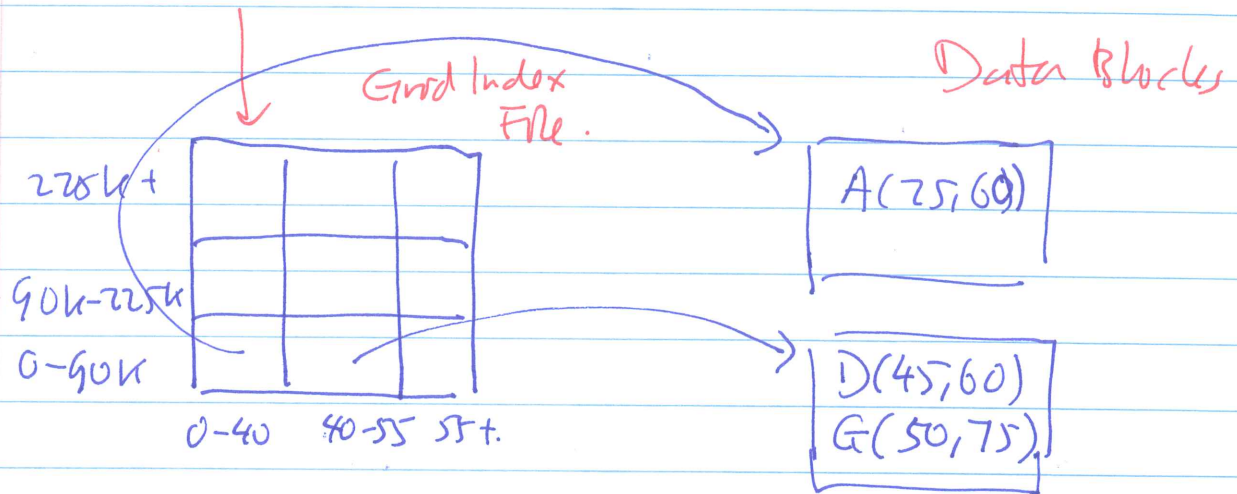
(age, salary).

- Data (input):

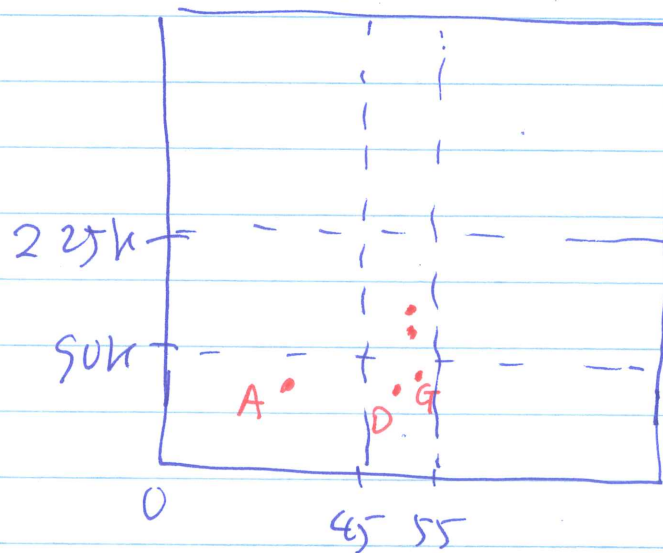
A(25, 6)	D(45, 60)	G(50, 75)	J(50, 100)
B(50, 120)	E(70, 110)	H(85, 140)	K(30, 260)
C(25, 400)	F(45, 350)	I(50, 275)	L(60, 260)

Ranges: Age: 0-40, 40-55, 55-~~+~~
Salary: 0-90k, 90k-225k, 225k+

Grid Index File for these records:



The book represent it as follows:



Lookup in a Grid Index file

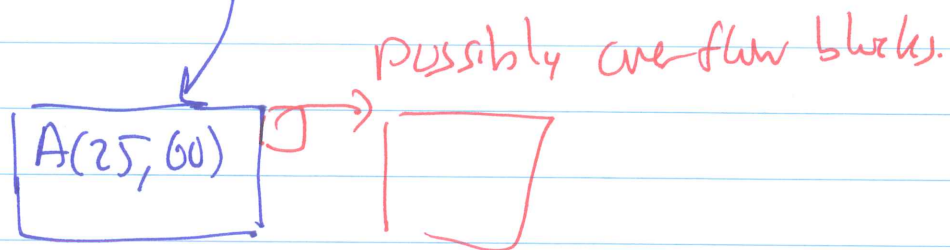
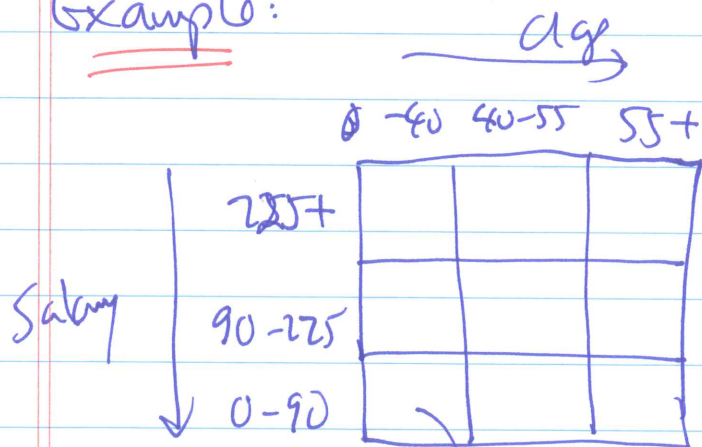
Recall: Grid Index file contains:

- (1) Number of dimensions
- (2) List of values at which the grid lines occur in each dimension.

Given: a "point" = value in N -dimension
 $(v_1, v_2, v_3, \dots, v_N)$.

$\Rightarrow$ determine the "position"
of the "point" in the grid
for each dimension.

Example:



Given point:

age = 25 → 0-40
Salary = 60 → 0-90

read table ptr.

Inserting into a Good Index File

(1) Do Lookup first

→ find blk pointer.

(2) If there is room in data block

⇒ insert

Overflow:

(A) Add overflow block ~~if~~

(B) Reorganize the good lines.

It's like "dynamic hashing"

but more complex.



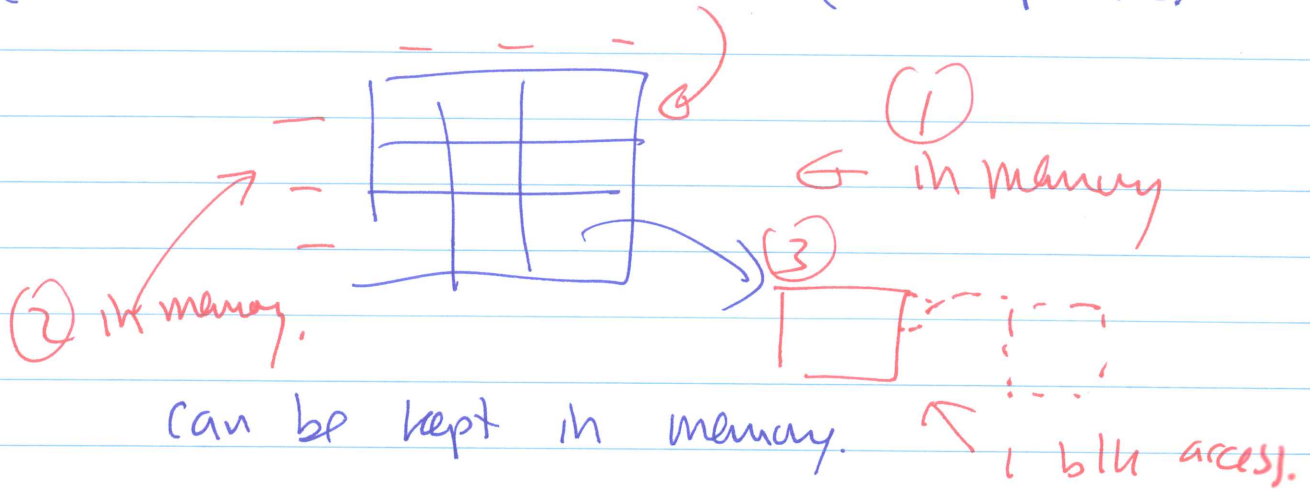
New good line will divide
MANY buckets !!

(could make many small/empty
buckets!).

Performance of Grid Index

Assumptions:

(1) The bucket matrix: (grid of blk pointers)



(2) Values of the grid lines

can also be kept in memory.

(3) Few overflow blocks.

(SU access to data block $\approx$ 1 blk)

Performance of Grid Index - Numbers.

(1) Look up a specific point $(v_1, v_2, \dots, v_N)$

Ø blk access to get blk pointer.

↓ blk access to get data block.

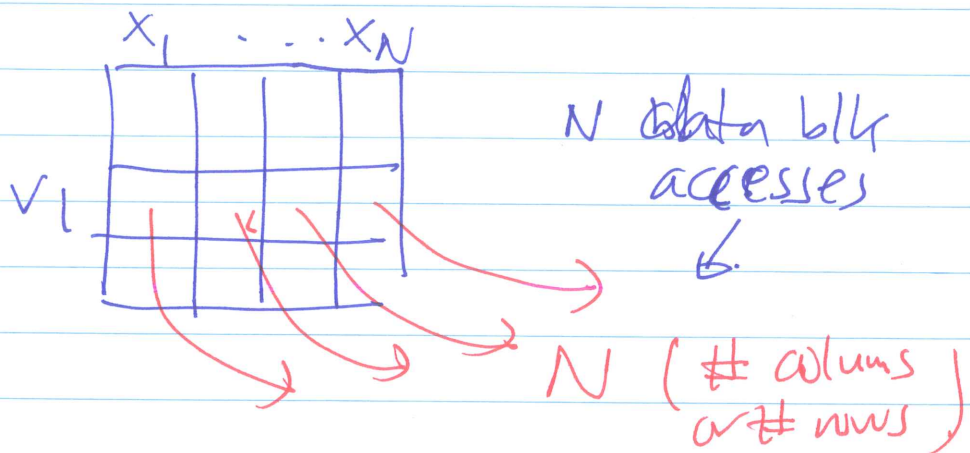
Insert/Delete:

may require an additional write blk.

(or overflow).

(2) Partial Match query: $(v_1, *)$

Ø blk access to get blk pointers

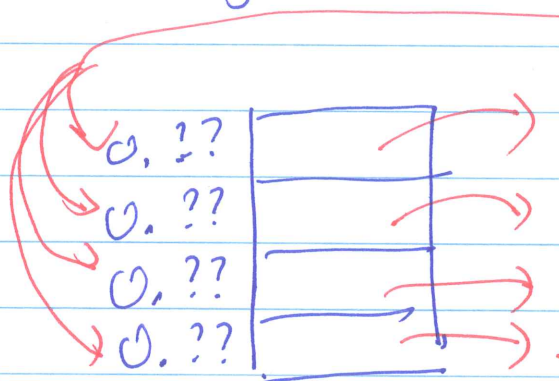


- Partial Match search: (example)

age = 50
 $\uparrow$
 given

Salary = *
 $\uparrow$
 unknown.

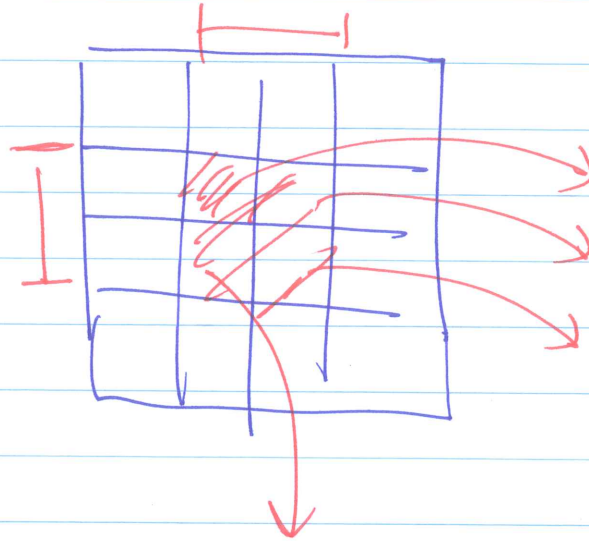
Records for age = 50 are found here:



- Look up specific point: very fast

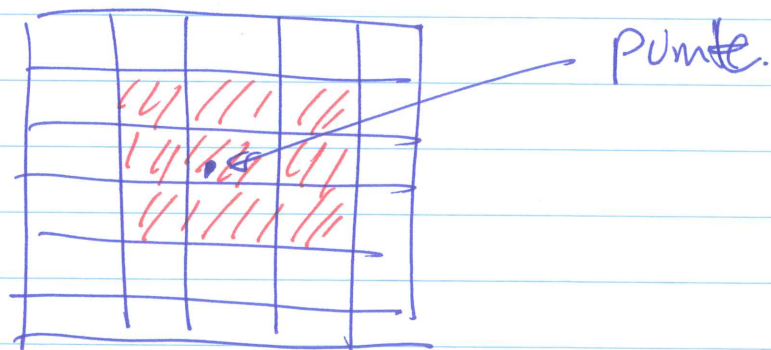
(it's hashing after all...)

③ Range Queries:



All data blocks in the range.

④ Nearest Neighbor Query:

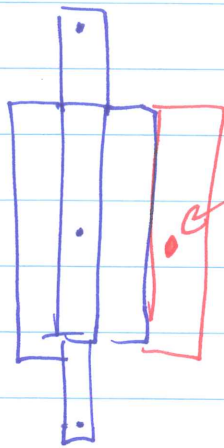


Answer: 9 data blocks

CAVEAT

It actually depends on the geometry

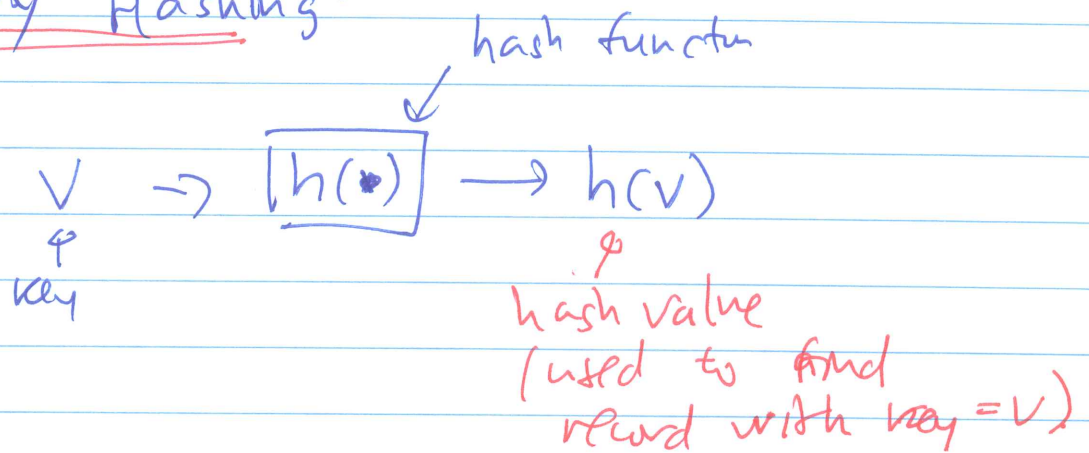
If the Grid is very stretched:



Nearest Neighbor!!!

Partitioned Hashing:

• Ordinary Hashing:



• Partitioned Hashing:

Key is composite: $(V_1, V_2, \dots, V_n)$.

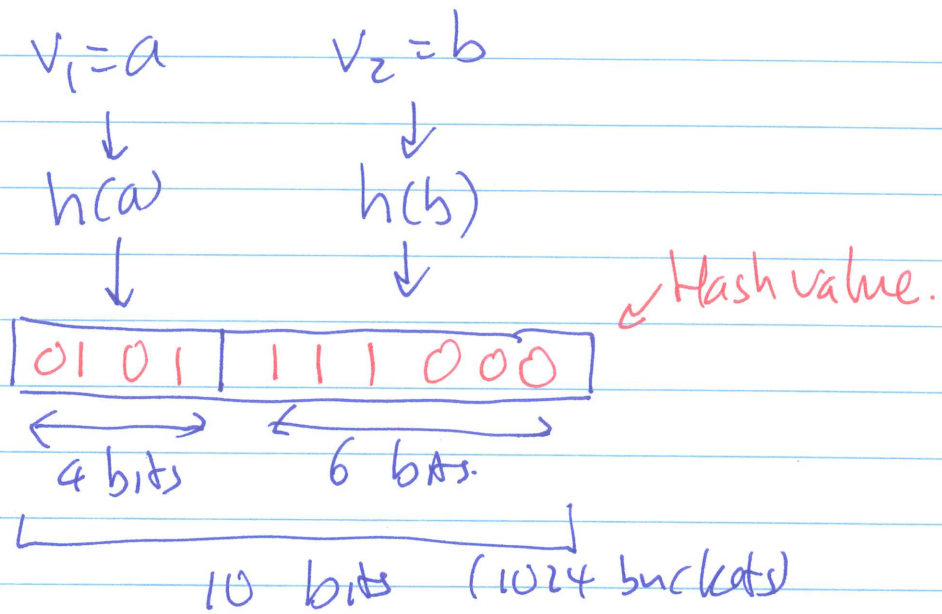
Uses n hash functions:

$$h_1(\cdot), h_2(\cdot), \dots, h_n(\cdot)$$

hash value = concatenation of

$$\underline{h_1(V_1) \cdot h_2(V_2) \cdot \dots \cdot h_n(V_n)}$$

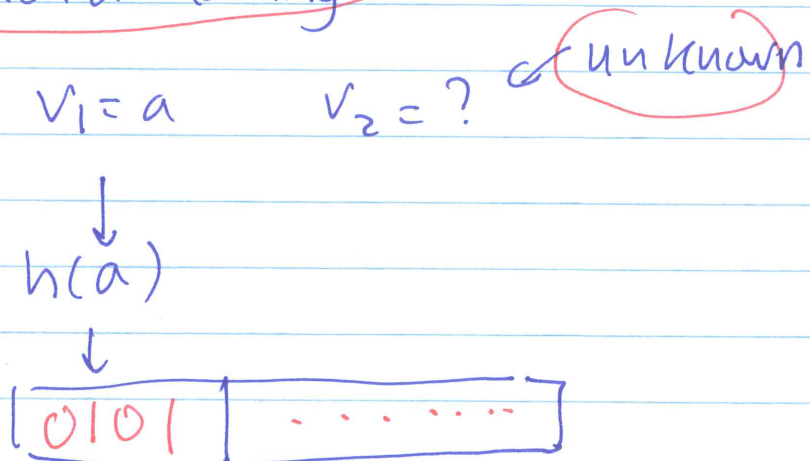
Example:



• Difference between Ordinary Hashing vs Partitioned Hashing

① Ordinary Hashing cannot hash when one (or more) keys are missing.

Partitioned Hashing:

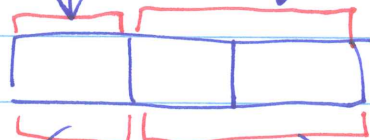


↳ you can still get a RANGE (subset) of buckets to search!

Concrete Example: who buys jewelry data

age = ... salary

$h_1(\cdot)$ $h_2(\cdot)$



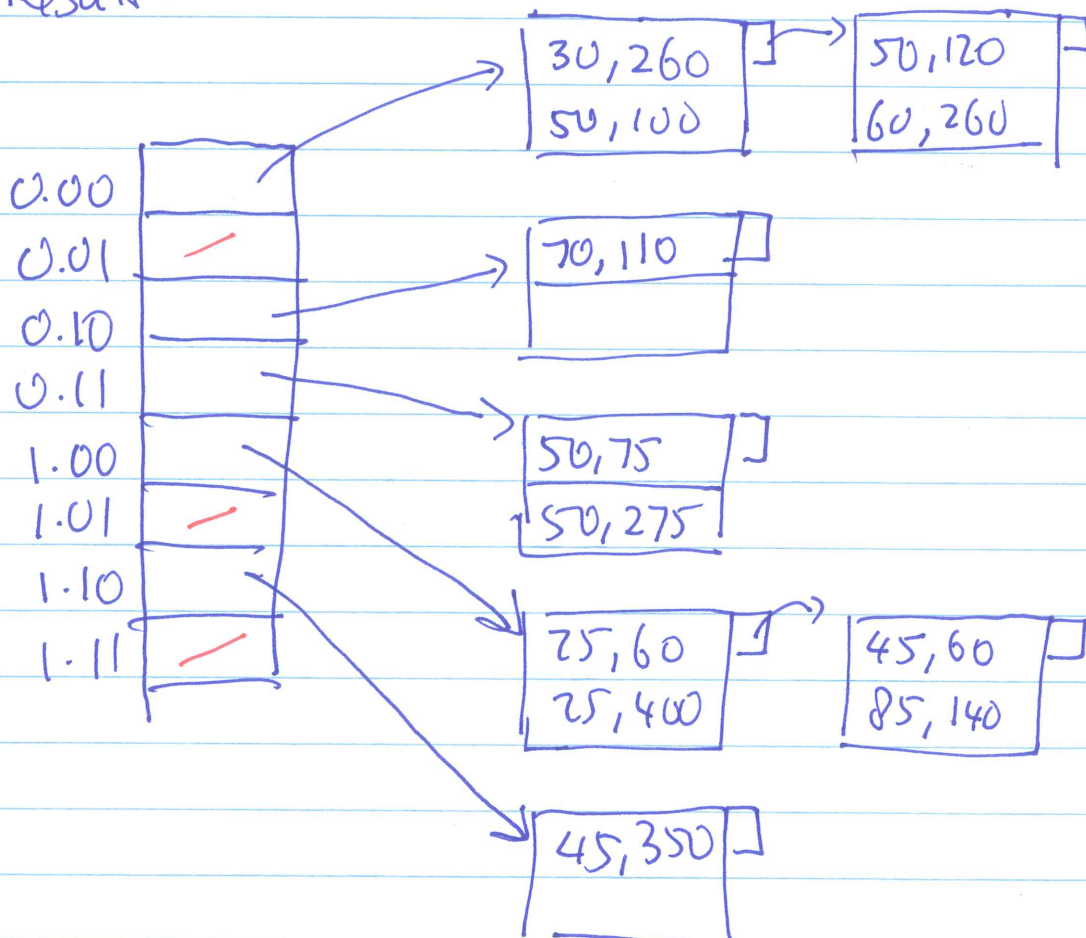
3 bit hash value

$$h_1(\cdot) = \text{age} \% 2$$

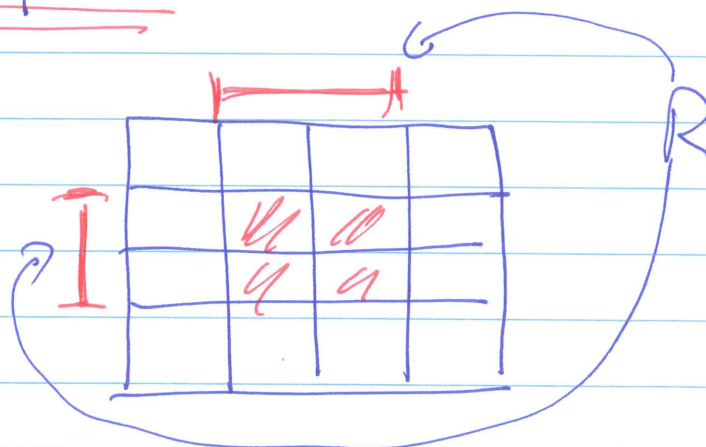
$$h_2(\cdot) = (\text{salary} / 1000) \% 4$$

Assume:
2 rewards
per block

Result:



- Range queries:



- Hash the values in the range $(v_1, v_2) \in R$
 $\rightarrow h(v_1, v_2)$
- Use the $h(v_1, v_2)$ to access the blocks.

- Nearest Neighbor:

Partitioned Hashing is

completely useless for Nearest Neighbor

(because there is no notion of distance in the hash function.

↓
closeness of the bucket numbers

≠ physical distance between data "points" (records).

Advantage of Partitioned Hashing.

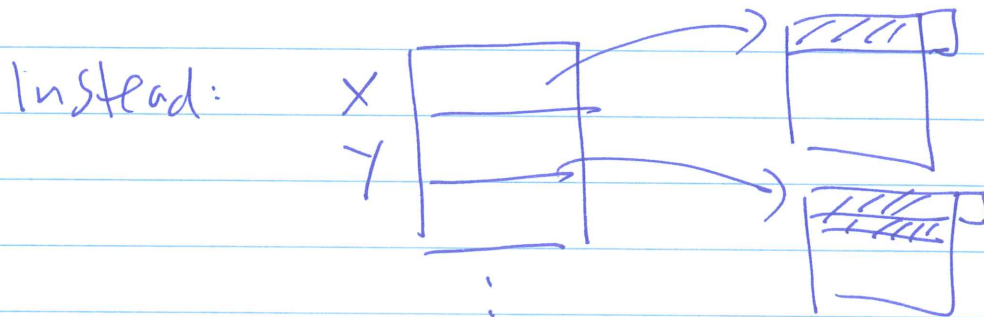
- Good hash function can randomize the records :-

Good occupation rate
(avg # records per bucket similar for all buckets.)

- Good files (especially with LARGE # dimensions) will have

many empty buckets.

Handling Tiny Buckets:



You can pack different buckets
into one block:

Block header must contain information
on:

- (1) bucket ID
- (2) records in specific bucket.

Eg:

