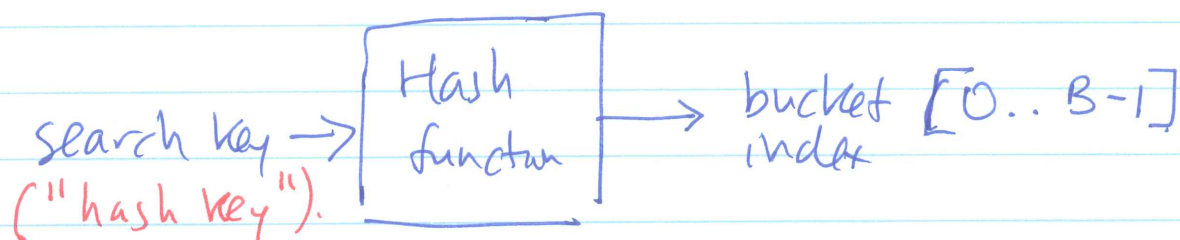


3

Hash Tables

• Hash Function:

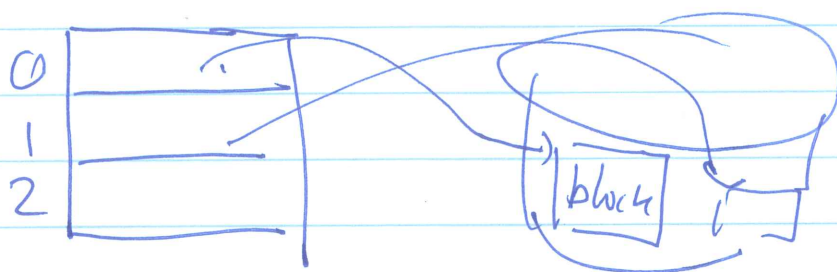


$K = \text{search key}$

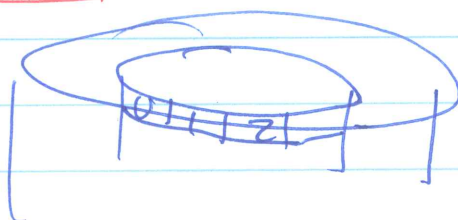
$h(K) = \text{bucket index}$

• Bucket Array:

(1) array of block pointers in memory



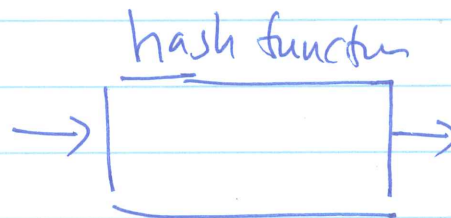
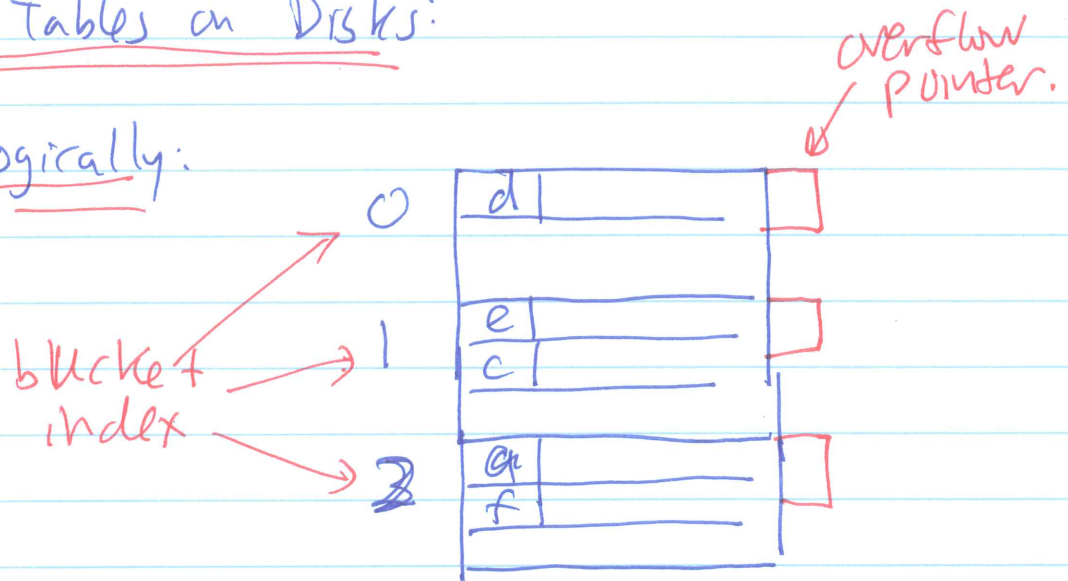
(2) Special (reserved) tracks/cylinders on disk



(each bucket index = 1 block pointer location)

Hash Tables on Disks:

Logically:

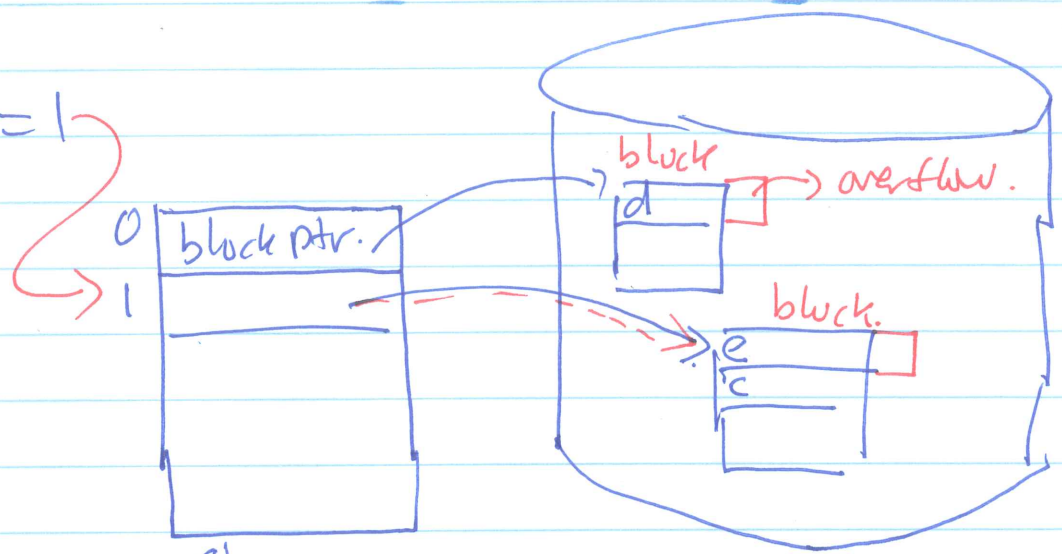


$$\begin{aligned}h(d) &= \phi \\h(e) &= 1 \\h(c) &= 1\end{aligned}$$

$$\begin{aligned}h(a) &= 2 \\h(f) &= 2\end{aligned}$$

Physically:

$h(e) = 1$



(can be :

(1) array of block pointers
in memory

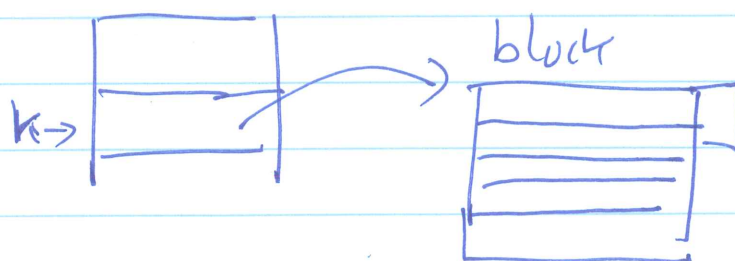
(2) special (reserved) cylinder
on disk.

Traditional Hash Index

- Insert (K):

(1) compute $k = h(K)$

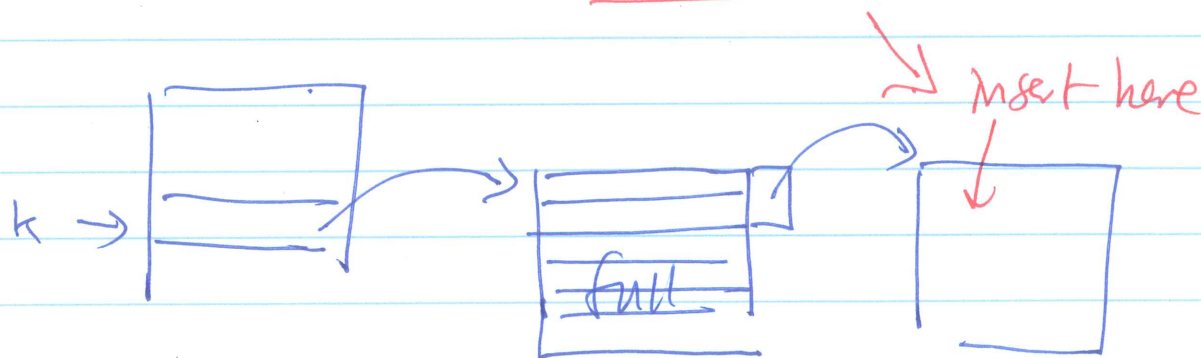
(2) Find hash bucket :



(3) If hash bucket has space
→ insert record.

(4) If hash bucket is full

→ allocate overflow block.

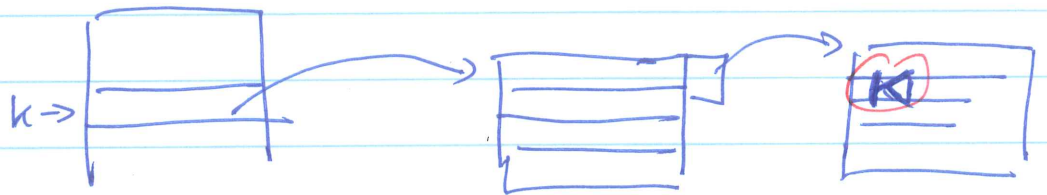


Delete from Hash Table.

• Delete (K)

(1) compute $k = h(K)$.

(2) Find hash bucket:



Delete record if found.

(3) Optionally:

consolidate blocks
at a bucket into
few blocks.

NAE:

consolidation can lead
to oscillation!!!

policy
(threshold)

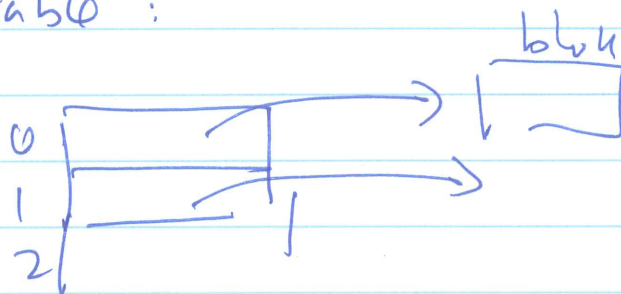
(insert → delete → insert
split / join / split)

Efficiency of Hash Table indexes.

- ideal: no overflow blocks.

⇒ 1 block access to find record.

(assuming that the block ptr table :



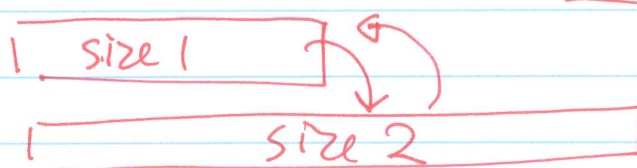
can fit in memory.

- With many overflow blocks

→ slower.

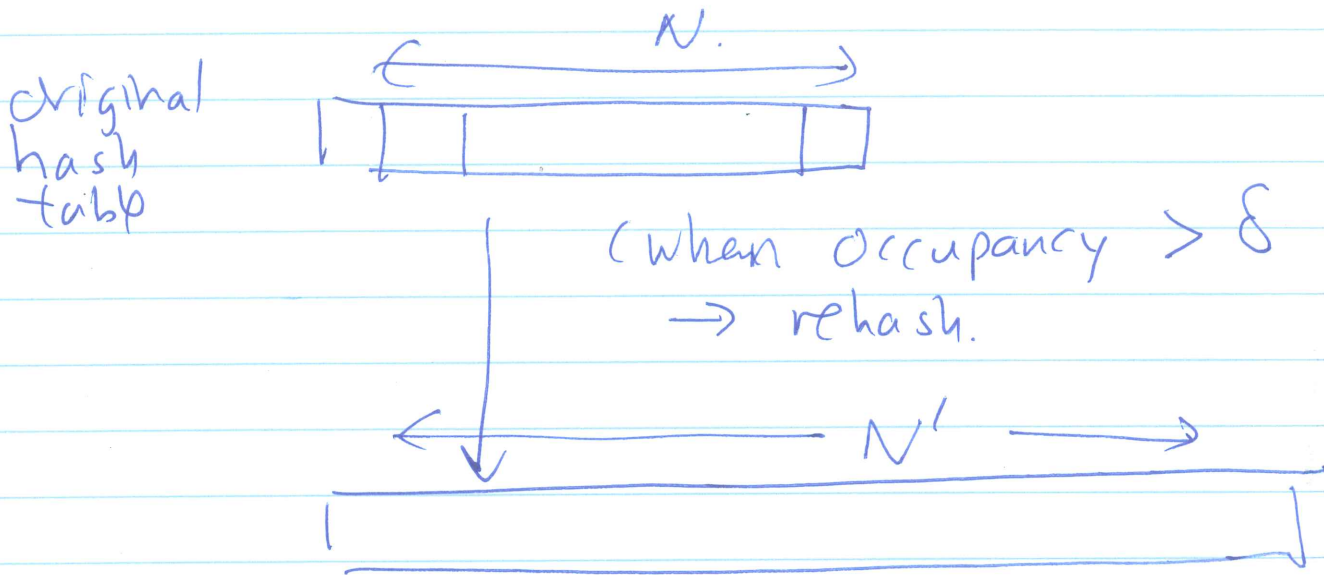
Solution

⇒ Dynamic size hash table.



Trivial Dynamic Sized Hash Table

= Re-hashing.



(1) Go through every record (key) in original hash table.

(2) Compute new hash index (location) and store in new hash table.

Problem: Very time consuming!