

RELATIONAL MODEL UNIT-III

Introduction to Relational Model, Basic Structure, Database Schema, Relational Algebra, Relational Calculus. File Organization and Indexing: Introduction, Types of File Organizations, Overview of Indexes, Types of Indexes, Index Data Structures, Tree structured Indexing, Hash based Indexing.

INTRODUCTION TO RELATIONAL MODEL

Relational data model is the most popular data model used widely around the world for data storage. In this model data is stored in the form of tables.

Relational Model Concepts

Table is also called Relation. Let the below table name be STUDENT_DATA

Attribute / Column / Field

Degree=No of columns=4

htno	Name	age	city
501	Amar	19	Hyderabad
502	Akbar	18	Warngal
503	Antony	19	Karimnagar

Tuple / Row / Record

Cardinality=No of rows=3

Table: In relational model the data is saved in the form of tables. A table has two properties rows and columns. Rows represent records and columns represent attributes.

Attribute: Each column in a Table is an attribute. Attributes are the properties that define a relation. e.g., HTNO, NAME, AGE, CITY in the above relation.

Tuple: Every single row of a table is called record or tuple.

Relation Schema: It represents the name of the relation (Table) with its attributes.

Eg., STUDENT_DATA(htno, name, age, city)

Degree: The total number of attributes in the relation is called the degree of the relation.

Cardinality: Total number of rows present in the Table.

RELATIONAL ALGEBRA

Relational Algebra is procedural query language, which takes Relation as input and generates relation as output. Relational algebra mainly provides theoretical foundation for relational databases and SQL.



Operator Symbol	Operator Name	Explanation
π	Projection	Select column names
σ	Selection	Select row values
ρ	Renaming	Rename a table name or expression results
$\cup$	Union	Perform union operation
$\cap$	Intersection	Perform inter section operation
$-$	Setdeference	Perform set difference operation
$\times$	Cartesianproduct	Every row of first table is joined with every row of second table
$\bowtie$	Join	Join two tables based on some condition

- i. **Select Operation(σ):** It select stuples that satisfy the given predicate from a relation.

Notation: $\sigma_p(r)$

where σ stands for selecting tuples (rows) and r stands for relation (table) name. p is propositional logic formula which may use connectors like **and**, **or**, and **not**. These terms may use relational operators like $=, \neq, \geq, <, >, \leq$.

Example1: $\sigma_{\text{subject}=\text{"database"}}(\text{Books})$

Output: Selects rowswhosesubjectis'database'frombooks table.

Example2: $\sigma_{\text{subject}=\text{"database"} \text{ and } \text{price}=\text{"450"}}(\text{Books})$

Output: Selects rowsfrombookswheresubjectis'database'and'price'is450.

Example3: $\sigma_{\text{subject}=\text{"database"} \text{ and } \text{price}=\text{"450"} \text{ or } \text{year} > \text{"2010"}}(\text{Books})$

Output: Selects rows from books where subject is'database'and'price'is450orthose books published after 2010.

- ii. **Project Operation (Π):**It projects column(s)thatsatisfyagiven predicate.

Notation: $\Pi_{A_1, A_2, \dots, A_n}(r)$

where A_1, A_2, A_n are column(attribute)namesofrelation r . Duplicaterowsare automatically eliminated in the output.

Example: $\Pi_{\text{subject}, \text{author}}(\text{Books})$

Displayvalues fromcolumnssubjectand authorfrom therelation Books.

- iii. **Union Operation ($\cup$):**It performs union operation between two given relations.Itcombines rows from two given relations.

Notation: $r \cup s$

Where r and s are either database relations or relation result set (temporary relation). $r \cup s$ returns a relation instance containing all tuples that occur in either relation instance r or relation instance s (or both). For a union operation to be valid, the following conditions must hold:

- Rands must have the same number of attributes.
- Attribute domains must be compatible in rands.

Example: $\Pi_{\text{author}}(\text{Books}) \cup \Pi_{\text{author}}(\text{Articles})$

Output: Projects the names of the authors who have either written a book or an article or both.

- iv. **Intersection Operation ($\cap$):** It performs intersection operation between two given relations. It collect only rows which are common in the two given relations.

Notation: $R \cap S$

$R \cap S$ returns a relation instance containing all tuples that occur in *both* R and S . The relations R and S must be union-compatible, and the schema of the result is defined to be identical to the schema of R .

$\Pi_{\text{author}}(\text{Books}) \cap \Pi_{\text{author}}(\text{Articles})$

Output: Projects the names of the authors who have written both book and an article.

- v. **Set Difference ($-$):** It finds tuples which are present in one relation but not in the second relation.

Notation: $r - s$

Finds all the tuples that are present in r but not in s .

Example: $\Pi_{\text{author}}(\text{Books}) - \Pi_{\text{author}}(\text{Articles})$

Output– Provides the name of authors who have written books but not articles.

- vi. **Cartesian Product ($\times$):** It returns a relation instance whose schema contains all the fields of table-1 (in the same order as they appear in table-1) followed by all the fields of table-2. It combines every row in first table with every row in the second table.

Notation: $r \times s$

Where r and s are relations and their output will be defined as $r \times s = \{ q \mid q \in r \text{ and } t \in s \}$

- vii. **Natural join ($\bowtie$):** The most general version of the join operation accepts a *join condition* C and a pair of relation instances as arguments and returns a relation instance. The *join condition* is identical to a *selection condition* in form. The operation is defined as follows:

$$R \bowtie S = \sigma_c(R \times S)$$

Thus $\bowtie$ is defined to be a cross-product followed by a selection. Note that the condition c can refer to attributes of both R and S .

Note: If the condition c in $R \bowtie S$ contain equal operator, then it is called equi-join

viii. **Natural Join ($\bowtie$):** In this case, we can simply omit the join condition; the default is that the join condition is a collection of equalities on all common fields. We call this special case as natural join, and it has the nice property that the result is guaranteed not to have two fields with the same name.

ix. **Rename Operation (ρ):** The results of relational algebra are also relations but without any name. The rename operation allows us to rename the output relation. 'rename' operation is denoted with small Greek letter ρ .

Notation: $\rho(\text{temp}, E)$

Where the result of expression E is saved with name of **temp**.

x. **Division ($/$):** Consider two relation instances A and B in which A has (exactly) two fields x and y and B has just one field y , with the same domain as in A . We define the *division* operation A / B as the set of all x values (in the form of unary tuples) such that for every y value in (a tuple of) B , there is a tuple (x, y) in A .

Example:

A	SNO	PNO	B1	PNO	A/B1	SNO
	S1	P1		P2		S1
	S1	P2				S2
	S1	P3	B2	PNO	A/B2	SNO
	S1	P4		P2		S1
	S2	P1		P4		S4
	S2	P2	B3	PNO	A/B3	SNO
	S3	P2		P2		S1
	S4	P2		P4		S4
	S4	P4				

Sample Queries: We present a number of sample queries using the following schema:

Sailors (sid: integer, sname: string, rating: integer, age: real)

Boats (bid: integer, bname: string, color: string) Reserves

(sid: integer, bid: integer, day: date)

The key fields are underlined, and the domain of each field is listed after the field name. Thus *sid* is the key for Sailors, *bid* is the key for Boats, and all three fields together form the key for Reserves. Fields in an instance of one of these relations will be referred to by name, or positionally, using the order in which they are listed above.

(Q1) Find the names of sailors who have reserved boat 103.

This query can be written as follows:

$\pi_{sname}((\sigma_{bid=103}Reserves) \bowtie Sailors)$

We first compute the set of tuples in Reserves with *bid* = 103 and then take the natural join of this set with Sailors. This expression can be evaluated on instances of Reserves and Sailors. Evaluated on the instances *R2* and *S3*, it yields a relation

(Q2) Find the names of sailors who have reserved a red boat.

$\pi_{sname}((\sigma_{color='red'}Boats) \bowtie Reserves \bowtie Sailors)$

This query involves a series of two joins. First we choose (tuples describing) red boats. Then, we join this set with Reserves (natural join, with equality specified on the **bid** column) to identify reservations of red boats. Next, we join the resulting intermediate relation with Sailors (natural join, with equality specified on the **sid** column) to retrieve the names of sailors who have made reservations for red boats. Finally, we project the sailors' names.

(Q3) Find the colors of boats reserved by Lubber.

$\pi_{color}((\sigma_{sname='Lubber'}Sailors) \bowtie Reserves \bowtie Boats)$

(Q4) Find the names of sailors who have reserved at least one boat.

$\pi_{sname}(Sailors \bowtie Reserves)$

The join of Sailors and Reserves creates an intermediate relation in which tuples consist of a Sailors tuple 'attached to' a Reserves tuple. A Sailors tuple appears in (some tuple of) this intermediate relation only if at least one Reserves tuple has the same *sid* value, that is, the sailor has made some reservation.

(Q5) Find the names of sailors who have reserved a red or a green boat.

$\rho(Tempboats, (\sigma_{color='red'}Boats) \cup (\sigma_{color='green'}Boats))$

$\pi_{sname}(Tempboats \bowtie Reserves \bowtie Sailors)$

We identify the set of all the rows that are either red or green from boats table. We rename this result as Tempboats. Then we join Temp boats with Reserves to identify sid's of sailors. Finally, we join with Sailors to find the names of Sailors with those sids.

(Q6) Find the names of sailors who have reserved a red and a green boat

$$\rho(\text{Tempboats2}, (\sigma_{\text{color}='red'}\text{Boats}) \cap (\sigma_{\text{color}='green'}\text{Boats})) \\ \pi_{\text{sname}}(\text{Tempboats2} \bowtie \text{Reserves} \bowtie \text{Sailors})$$

However, this solution is incorrect—it instead tries to compute sailors who have reserved a boat that is both red and green. A boat can be only one color; this query will always return an empty answer set. The right answer is

$$\rho(\text{Tempred}, \pi_{\text{sid}}((\sigma_{\text{color}='red'}\text{Boats}) \bowtie \text{Reserves})) \\ \rho(\text{Tempgreen}, \pi_{\text{sid}}((\sigma_{\text{color}='green'}\text{Boats}) \bowtie \text{Reserves})) \\ \pi_{\text{sname}}((\text{Tempred} \cap \text{Tempgreen}) \bowtie \text{Sailors})$$

The two temporary relations compute the sids of sailors, and their intersection identifies sailors who have reserved both red and green boats.

(Q7) Find the names of sailors who have reserved at least two boats.

$$\rho(\text{Reservations}, \pi_{\text{sid}, \text{sname}, \text{bid}}(\text{Sailors} \bowtie \text{Reserves})) \\ \rho(\text{Reservationpairs}(1 \rightarrow \text{sid1}, 2 \rightarrow \text{sname1}, 3 \rightarrow \text{bid1}, 4 \rightarrow \\ \text{sid2}, 5 \rightarrow \text{sname2}, 6 \rightarrow \text{bid2}), \text{Reservations} \times \text{Reservations}) \\ \pi_{\text{sname1}}(\sigma_{(\text{sid1}=\text{sid2}) \cap (\text{bid1} \neq \text{bid2})} \text{Reservationpairs})$$

First, we compute tuples of the form (sid, sname, bid), where sailor sid has made a reservation for boat bid; this set of tuples is the temporary relation Reservations. Next we find all pairs of Reservations tuples where the same sailor has made both reservations and the boats involved are distinct. Here is the central idea: To show that a sailor has reserved two boats, we must find two Reservations tuples involving the same sailor but distinct boats. Finally, we project the names of such sailors.

(Q8) Find the sids of sailors with age over 20 who have not reserved a red boat.

$$\pi_{\text{sid}}(\sigma_{\text{age} > 20} \text{Sailors}) - \pi_{\text{sid}}((\sigma_{\text{color}='red'}\text{Boats}) \bowtie \text{Reserves} \bowtie \text{Sailors})$$

This query illustrates the use of the set-difference operator. Again, we use the fact that sid is

The key for Sailors. We first identify sailors aged over 20 instances and then discard those who have reserved a red boat to obtain the answer.

(Q9) Find the names of sailors who have reserved all boats.

The use of the word *all* (or *every*) is a good indication that the division operation might be applicable:

$$\rho(\text{Tempsids}, (\pi_{\text{sid}, \text{bid}} \text{Reserves}) / (\pi_{\text{bid}} \text{Boats}))$$

$$\pi_{\text{sname}}(\text{Tempsids} \bowtie \text{Sailors})$$

(Q10) Find the names of sailors who have reserved all boats called Interlake.

$$\rho(\text{Tempsids}, (\pi_{\text{sid}, \text{bid}} \text{Reserves}) / (\pi_{\text{bid}} (\sigma_{\text{bname}='Interlake'} \text{Boats})))$$

$$\pi_{\text{sname}}(\text{Tempsids} \bowtie \text{Sailors})$$

RELATIONAL CALCULUS

Relational calculus is an alternative to relational algebra. In contrast to the algebra, which is procedural, the calculus is nonprocedural, or *declarative*, in that it allows us to describe the set of answers without being explicit about how they should be computed.

Tuple Relational Calculus

Tuple Relational Calculus is a **non-procedural query language** unlike relational algebra. Tuple Calculus provides only the description of the query but it does not provide the methods to solve it. Thus, it explains what to do but not how to do.

In Tuple Relational Calculus, a query is expressed as $\{t \mid P(t)\}$

Where t = resulting tuples, $P(t)$ = known as Predicate and these are the conditions that are used to fetch t . Thus, it generates set of all tuples t , such that Predicate $P(t)$ is true for t .

$P(t)$ may have various conditions logically combined with OR($\vee$), AND($\wedge$), NOT($\neg$). It also uses quantifiers:

$\exists t \in r(Q(t))$ = "there exists" a tuple t in relation r such that predicate $Q(t)$ is true.

$\forall t \in r(Q(t))$ = $Q(t)$ is true "for all" tuples in relation r .

(Q12) Find the names and ages of sailors with arating above7.

$$\{P \mid \exists S \in \text{Sailors} (S.\text{rating} > 7 \wedge P.\text{name} = S.\text{sname} \wedge P.\text{age} = S.\text{age})\}$$

This query illustrates a useful convention: P is considered to be a tuple variable with exactly two fields, which are called *name* and *age*.

(Q13) Find the sailername, boatid, and reservation date for each reservation

$$\{P \mid \exists R \in \text{Reserves} \quad \exists S \in \text{Sailors} \\ (R.\text{sid} = S.\text{sid} \wedge P.\text{bid} = R.\text{bid} \wedge P.\text{day} = R.\text{day} \wedge P.\text{sname} = S.\text{sname})\}$$

(Q1) Find the names of sailors who have reserved boat103. (similarquestionQ1from relational algebra)

$$\{P \mid \exists S \in \text{Sailors} \exists R \in \text{Reserves} (R.\text{sid} = S.\text{sid} \wedge R.\text{bid} = 103 \wedge P.\text{sname} = S.\text{sname})\}$$

This query can be read as follows: “Retrieve all sailor tuples for which there exists a tuple in Reserves, having the same value in the *sid* field, and with *bid* = 103.”

(Q2) Find the names of sailors who have reserved a red boat. (similarquestionQ2from relational algebra)

$$\{P \mid \exists S \in \text{Sailors} \exists R \in \text{Reserves} (R.\text{sid} = S.\text{sid} \wedge P.\text{sname} = S.\text{sname} \\ \wedge \exists B \in \text{Boats} (B.\text{bid} = R.\text{bid} \wedge B.\text{color} = 'red'))\}$$

This query can be read as follows: “Retrieve all sailor tuples S for which there exist tuples R in Reserves and B in Boats such that $S.\text{sid} = R.\text{sid}$, $R.\text{bid} = B.\text{bid}$, and $B.\text{color} = 'red'$.”

(Q7) Find the names of sailors who have reserved atleast twoboats. (similarquestionQ7 from relational algebra)

$$\{P \mid \exists S \in \text{Sailors} \exists R1 \in \text{Reserves} \exists R2 \in \text{Reserves} (S.\text{sid} = R1.\text{sid} \\ \wedge R1.\text{sid} = R2.\text{sid} \wedge R1.\text{bid} \neq R2.\text{bid} \wedge P.\text{sname} = S.\text{sname})\}$$

(Q9) Find the names of sailors who have reserved all boats. (similarquestionQ9from relational algebra)

$$\{P \mid \exists S \in \text{Sailors} \quad \forall B \in \text{Boats} \\ (\exists R \in \text{Reserves} (S.\text{sid} = R.\text{sid} \wedge R.\text{bid} = B.\text{bid} \wedge P.\text{sname} = S.\text{sname}))\}$$

(Q14) Find sailors who have reserved all red boats.

$$\{S \mid \exists S \text{ Sailors } \forall B \in \text{Boats}$$

$$(B.\text{color} = \text{'red'} \Rightarrow (\exists R \in \text{Reserves} (S.\text{sid} = R.\text{sid} \wedge R.\text{bid} = B.\text{bid})))\}$$

Domain Relational Calculus

A domain variable is a variable that ranges over the values in the domain of some attribute (e.g., the variable can be assigned an integer if it appears in an attribute whose domain is the set of integers).

ADRC query has the form $\{ \langle x_1, x_2, \dots, x_n \rangle \mid p(\langle x_1, x_2, \dots, x_n \rangle) \}$

where each x_i is either a *domain variable* or a constant and $p(\langle x_1, x_2, \dots, x_n \rangle)$ denotes a DRC formula whose only free variables are the variables among the x_i , $1 \leq i \leq n$. The result of this query is the set of all tuples $\langle x_1, x_2, \dots, x_n \rangle$ for which the formula evaluates to true.

A DRC formula is defined in a manner very similar to the definition of a TRC formula. The main difference is that the variables are now domain variables. Let op denote an operator in the set $\{<, >, =, \leq, \geq, \neq\}$ and let X and Y be domain variables. An atomic formula in DRC is one of the following:

- $(x_1, x_2, \dots, x_n) \in \text{Rel}$, where Rel is a relation with n attributes; each x_i , $1 \leq i \leq n$ is either a variable or a constant
- $X \text{ op } Y$
- $X \text{ op constant}$, or $\text{constant op } X$

A formula is recursively defined to be one of the following, where P and q are themselves formulas and $p(X)$ denotes a formula in which the variable X appears:

- any atomic formula
- $\neg p, P \wedge q, P \vee q, \text{ or } p \Rightarrow q$
- $\exists X(p(X))$, where X is a domain variable
- $\forall X(p(X))$, where X is a domain variable

(Q1) Find the names of sailors who have reserved boat 103.

$$\{ (N) \mid \exists I, T, A (\langle I, N, T, A \rangle \in \text{Sailors}$$

$$\wedge \exists Ir, Br, D (\langle Ir, Br, D \rangle \in \text{Reserves} \wedge Ir = I \wedge Br = 103)) \}$$

(Q2) Find the names of sailors who have reserved a red boat.

$$\{ \langle N \rangle \mid \exists I, T, A (\langle I, N, T, A \rangle \in \text{Sailors} \wedge \exists \langle I, Br, D \rangle \in \text{Reserves} \wedge \exists \langle Br, BN, 'red' \rangle \in \text{Boats}) \}$$

(Q7) Find the names of sailors who have reserved at least two boats.

$$\{ \langle N \rangle \mid \exists I, T, A (\langle I, N, T, A \rangle \in \text{Sailors} \wedge \exists Br1, Br2, D1, D2 (\langle I, Br1, D1 \rangle \in \text{Reserves} \wedge \langle I, Br2, D2 \rangle \in \text{Reserves} \wedge Br1 \neq Br2) \}$$

(Q9) Find the names of sailors who have reserved all boats.

$$\{ \langle N \rangle \mid \exists I, T, A (\langle I, N, T, A \rangle \in \text{Sailors} \wedge \forall B, BN, C (\neg (\langle B, BN, C \rangle \in \text{Boats}) \vee (\exists \langle Ir, Br, D \rangle \in \text{Reserves} (I = Ir \wedge Br = B)))) \}$$

File Organization and Indexing:

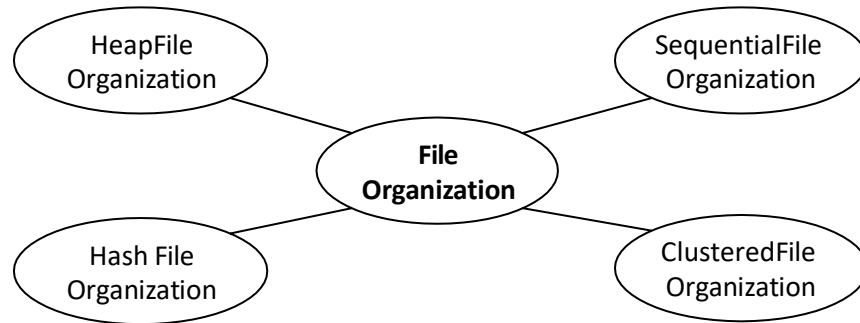
FILE ORGANIZATION

The database is stored as a collection of files. Each file contains a set of records. Each record is a collection of fields. For example, a student table (or file) contains many records and each record belongs to one student with fields (attributes) such as Name, Date of birth, class, department, address, etc.

File organization defines how file records are mapped onto disk blocks.

The records of a file are stored in the disk blocks because a block is the unit of data transfer between disk and memory. When the block size is larger than the record size, each block will contain more than one record. Sometimes, some of the files may have large records that cannot fit in one block. In this case, we can store part of a record on one block and the rest on another. A pointer at the end of the first block points to the block containing the remainder of the record.

The different types of file organization are given below:



Heap File Organization: When a file is created using Heap File Organization mechanism, the records are stored in the file in the order in which they are inserted. So the new records are inserted at the end of the file. In this type of organization inserting new records is more efficient. It uses linear search to search records.

Sequential File Organization: When a file is created using Sequential File Organization mechanism, all the records are ordered (sorted) as per the primary key value and placed in the file. In this type of organization inserting new records is more difficult because the records need to be sorted after inserting every new record. It uses binary search to search records.

Hash File Organization: When a file is created using Hash File Organization mechanism, a hash function is applied on some field of the records to calculate hash value. Based on the hash value, the corresponding record is placed in the file.

Clustered File Organization: Clustered file organization is not considered good for large databases. In this mechanism, related records from one or more relations are kept in a same disk block, that is, the ordering of records is not based on primary key or search key.

OVER VIEW OF INDEXES:

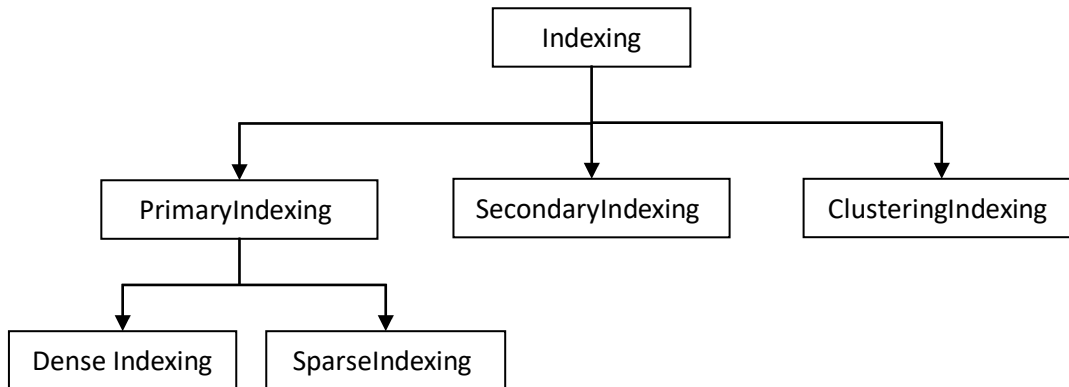
INDEXING

If the records in the file are in sorted order, then searching will become very fast. But, in most of the cases, records are placed in the file in the order in which they are inserted, so new records are inserted at the end of the file. It indicates, the records are not in sorted order. In order to make searching faster in the files with unsorted records, indexing is used.

Indexing is a data structure technique which allows you to quickly retrieve records from a database file. An Index is a small table having only two columns. The first column contains a copy of the primary or candidate key of a table. The second column contains a set of disk block addresses

where the record with that specific key value stored.

Indexing in DBMS can be of the following types:



i. Primary Index

- If the index is created by using the primary key of the table, then it is known as primary indexing.
- As primary keys are unique and are stored in a sorted manner, the performance of the searching operation is quite efficient.
- The primary index can be classified into two types: dense index and sparse index.

Dense index

- If every record in the table has one index entry in the index table, then it is called dense index.
- In this, the number of records (rows) in the index table is same as the number of records (rows) in the main table.
- As every record has one index entry, searching becomes faster.

TS	●	→	TS	Hyderabad	KCR
AP	●	→	AP	Amaravathi	Jagan
TN	●	→	TN	Madras	PalaniSwamy
MH	●	→	MH	Bombay	Thackray

Sparse index

- If only few records in the table have index entries in the index table, then it is called sparse index.
- In this, the number of records (rows) in the index table is less than the number of records (rows) in the main table.
- As not all the records have index entries, searching becomes slow for records that do not have

index entries.

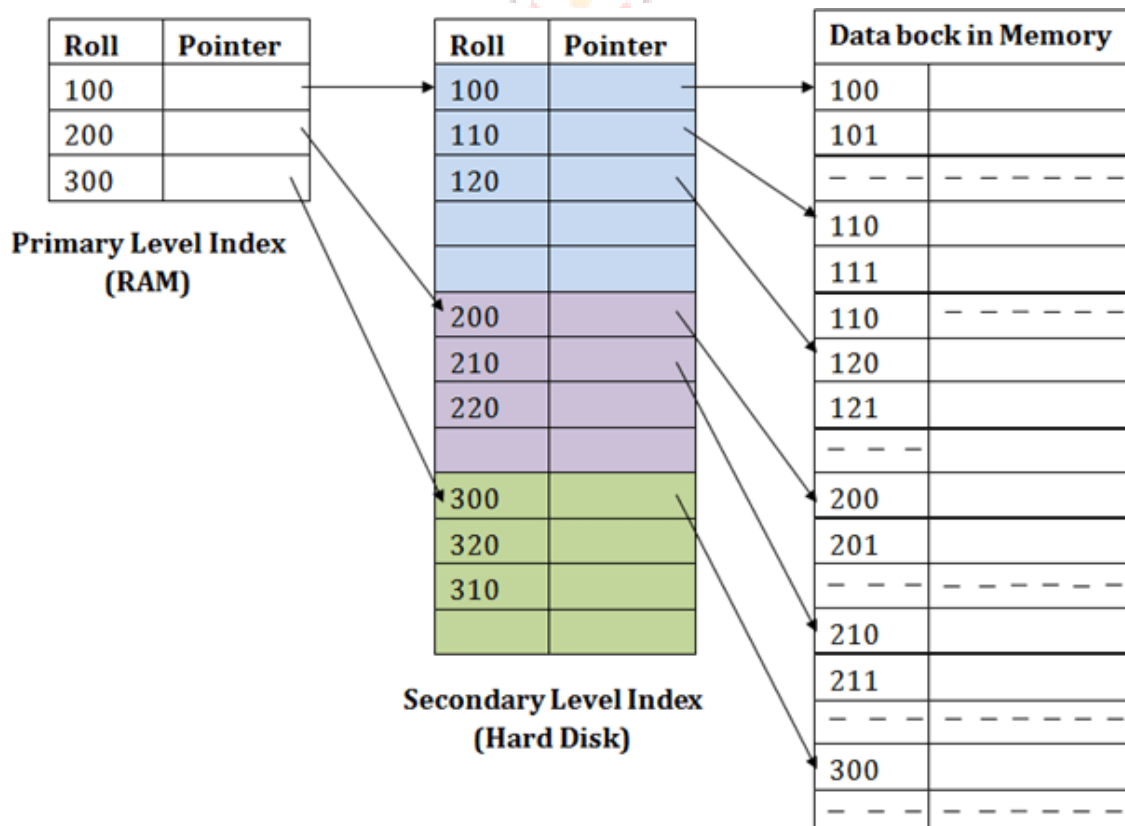


TS	●	→	TS	Hyderabad	KCR
TN	●	→	AP	Amaravathi	Jagan
MH	●	→	TN	Madras	PalaniSwamy
		→	MH	Bombay	Thackray

ii. Secondary Index

When the size of the main table grows, then the size of the index table also grows. If the index table size grows then fetching the address itself becomes slower. To overcome this problem, secondary indexing is introduced.

- In secondary indexing, to reduce the size of mapping, another level of indexing is introduced.
- It contains two levels. In the first level, each record in the main table has one entry in the first-level index table.
- The index entries in the first-level index table are divided into different groups. For each group, one index entry is created and added in the second-level index table.

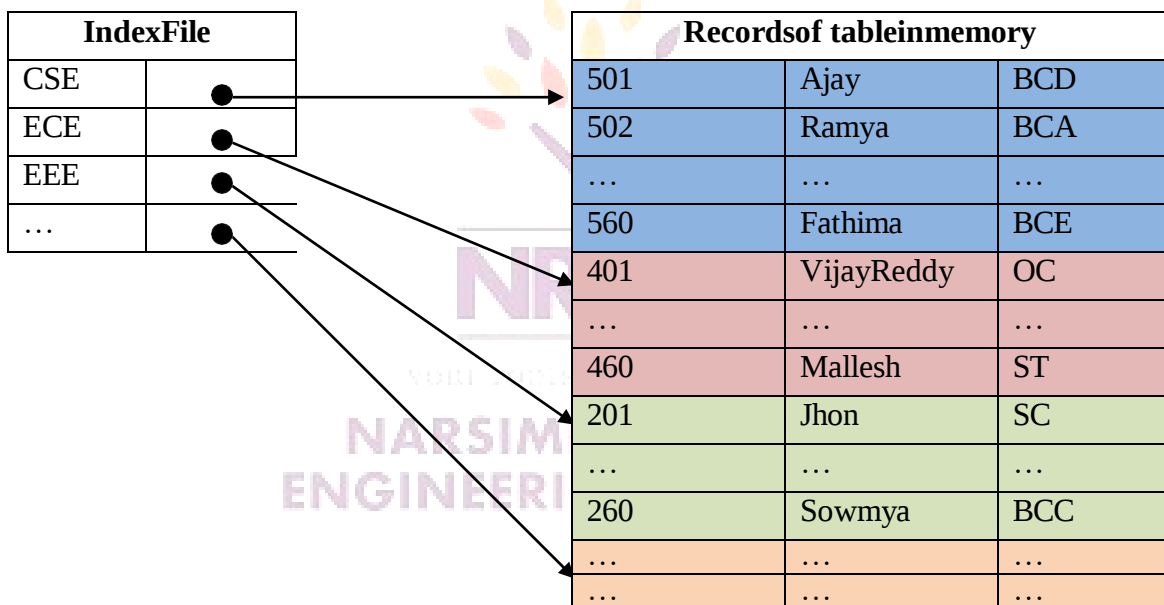


Multi-level Index: When the main table size becomes too large, creating a secondary level index improves the search process. Even if the search process is slow; we can add one more level of indexing and so on. This type of indexing is called a multi-level index.

iii. Clustering Index

- Sometimes the index is created on non-primary key columns which may not be unique for each record.
- In this case, to identify their record faster, we will group two or more columns to get the unique value and create index out of them. This method is called a clustering index.
- The records which have similar characteristics are grouped, and indexes are created for these group.

Example: Consider a college contains many students in each department. All the students belong to the same Dept_ID are grouped together and treated as a single cluster. One index pointers point to the one cluster as a whole. The index pointer points to the first record in each cluster. Here Dept_ID is a non-unique key.

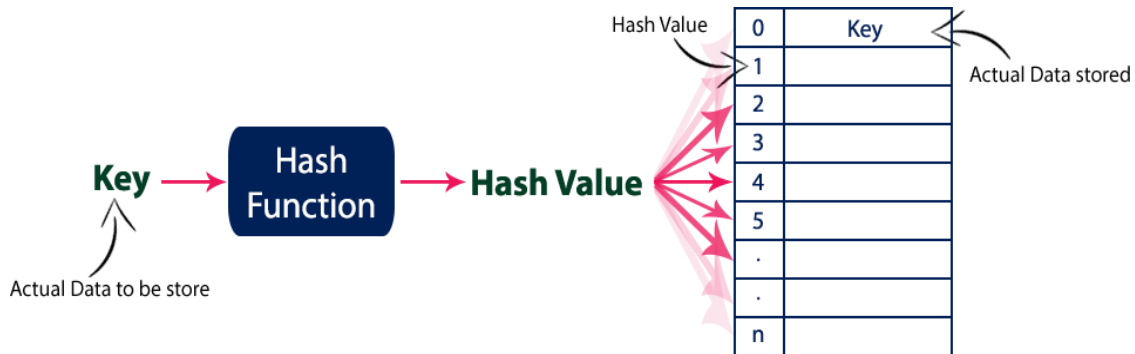


In above diagram we can see that, indexes are created for each department in the index file. In the data block, the students of each department are grouped together to form the cluster. The address in the index file points to the beginning of each cluster.

2. HASHBASED INDEXING

Hashing is a technique to directly search the location of desired data on the disk without using index structure. Hash function is a function which takes a piece of data (key) as input and produces a hash value as output which maps the data to a particular location in the hash table.

The concept of hashing and hash table is shown in the below figure



There are mainly two types of hashing methods:

- i. Static Hashing
- ii. Dynamic Hashing
 - Extended hashing
 - Linear hashing

3. STATIC HASHING

In static hashing, the hash function produces only fixed number of hash values. For example, consider the hash function

$$f(x) = x \bmod 7$$

For any value of x , the above function produces one of the hash value from $\{0, 1, 2, 3, 4, 5, 6\}$. It means static hashing maps search-key values to a fixed set of bucket addresses.

Example: Inserting 10, 21, 16 and 12 in hash table.

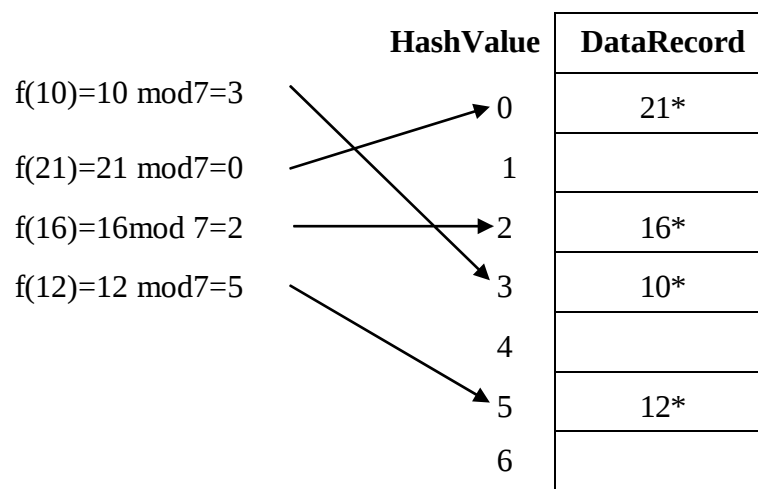


Figure 5.1: Static hashing

Suppose, latter if we want to insert 23, it produce hash value as 2 ($23 \bmod 7 = 2$). But, in the above hash table, the location with hash value 2 is not empty (it contains 16*). So, a collision occurs. To resolve this collision, the following techniques are used.

- Open addressing
- SeparateChainingorClosed addressing

i. Open Addressing:

Openaddressingisacollisionresolvingtechniquewhichstoresallthekeysinsidethe hash table. No key is stored outside the hash table. Techniques used for open addressing are:

- LinearProbing
- QuadraticProbing
- DoubleHashing

➤ LinearProbing:

In linear probing, when there is a collision, we scan forwards for the next empty slot to fill the key's record. If you reach last slot, then start from beginning.

Example: Consider figure 1. When we try to insert 23, its hash value is 2. But the slot with 2 is not empty. Then move to next slot (hash value 3), even it is also full, then move once again to next slot with hash value 4. As it is empty store 23 there. This is shown in the below diagram.

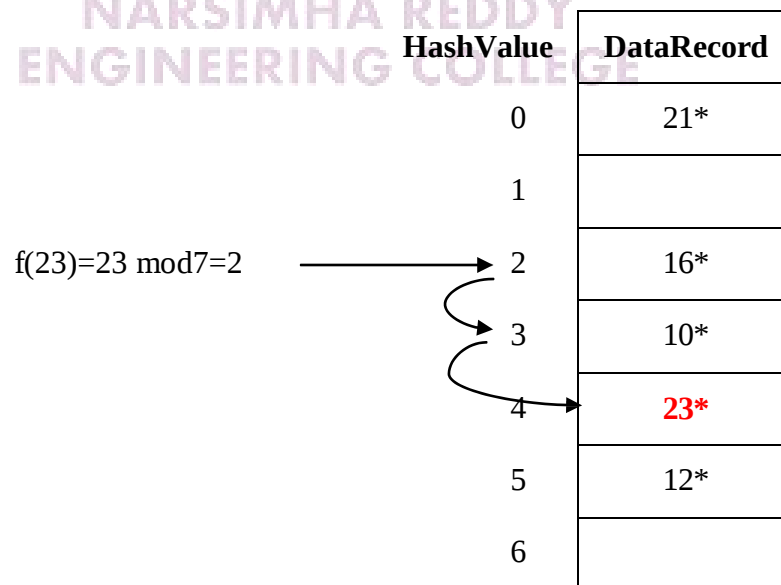


Figure5.2: Linear Probing

➤ Quadratic Probing:

In quadratic probing, when collision occurs, it computer new hash value by taking the original hash value and adding successive values of quadratic polynomial until an open slot is found. If here is a collision, it use the following hash function: $h(x) = (f(x) + i^2) \bmod n$, where $i = 1, 2, 3, 4, \dots$ and $f(x)$ is initial hash value.

Example: Consider figure 1. When we try to insert 23, its hash value is 2. But the slot with hash value 2 is not empty. Then compute new hash value as $(2 + 1^2) \bmod 7 = 3$, even it is also full, and then once again compute new hash value as $(2 + 2^2) \bmod 7 = 6$. As it is empty store 23 there. This is shown in the below diagram.

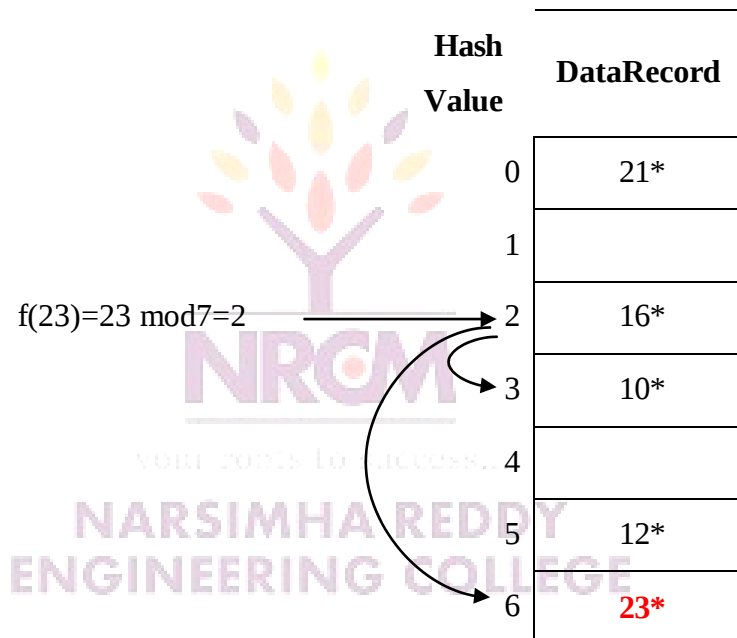


Figure5.3: Quadratic Probing

➤ DoubleHashing

In double hashing, there are two hash functions. The second hash function is used to provide an offset value in case the first function causes a collision. The following function is an example of double hashing: $(\text{firstHash}(\text{key}) + i * \text{secondHash}(\text{key})) \% \text{tableSize}$. Use $i = 1, 2, 3, \dots$

A popular second hash function is : $\text{secondHash}(\text{key}) = \text{PRIME} - (\text{key} \% \text{PRIME})$
 where PRIME is a prime smaller than the TABLE_SIZE.

Example: Consider figure1. When we try to insert 23, its hash value is 2. But the slot with hash value 2 is not empty. Then compute double hashing value as

$$\text{secondHash (key)} = \text{PRIME} - (\text{key} \% \text{PRIME}) \rightarrow \text{secondHash (23)} = 5 - (23 \% 5) = 2$$

$$\text{Doublehashing: (firstHash(key) + i * \text{secondHash(key)}) \% \text{tableSize} \rightarrow (2 + 1 * 2) \% 7 = 4$$

As the slot with hash value 4 is empty, store 23 there. This is shown in the below diagram.

	HashValue	DataRecord
	0	21*
	1	
$f(23) = 23 \bmod 7 = 2$	2	16*
	3	10*
	4	23*
	5	12*
	6	

Figure 5.4: Double Probing

ii. Separate Chaining or Closed addressing:

To handle the collision, This technique creates a linked list to the slot for which collision occurs. The new key is then inserted in the linked list. These linked lists to the slots appear like chains. So, this technique is called as *separate chaining*. It is also called as closed addressing.

Example: Inserting 10, 21, 16, 12, 23, 19, 28, 30 in hash table.

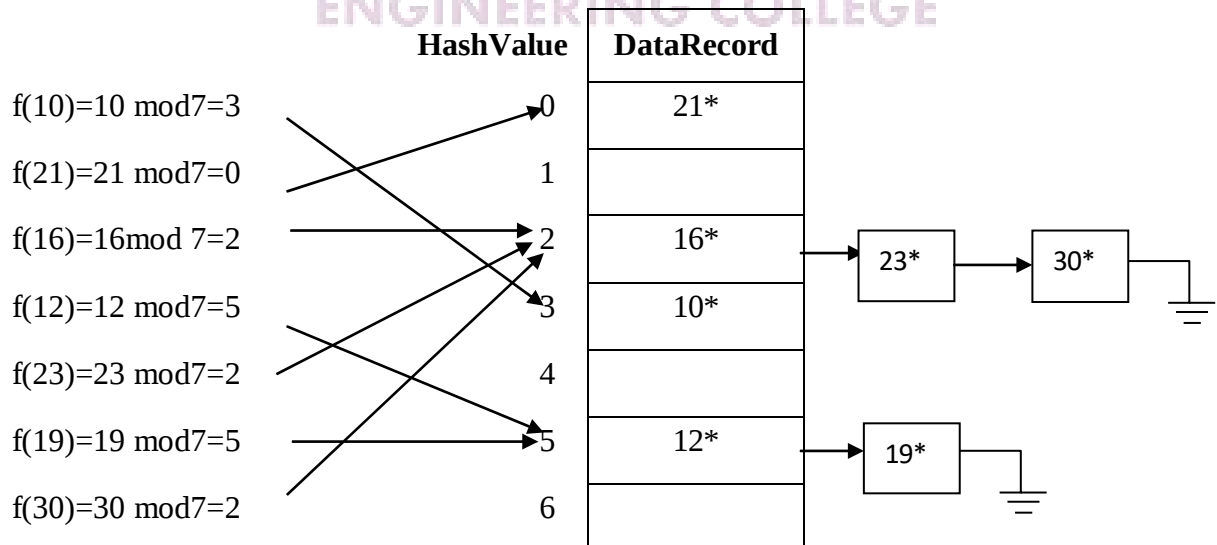


Figure 5.5: Separate chaining example

4. DYNAMIC HASHING

The problem with static hashing is that it does not expand or shrink dynamically as the size of the database grows or shrinks. Dynamic hashing provides a mechanism in which data buckets are added and removed dynamically and on-demand. Dynamic hashing can be implemented using two techniques. They are:

- Extended hashing
- Linear Hashing

i. Extendable hashing

In extendable hashing, a separate *directory* of pointers to buckets is used. The number of bits used in the directory is called *global depth* (gd) and number of entries in the directory = 2^{gd} . Number of bits used for locating the record in the buckets is called *local depth* (ld) and each bucket can store up to 2^{ld} entries. The hash function uses last few binary bits of the key to find the bucket. If a bucket overflows, it splits, and if local depth is greater than global depth, then the table doubles in size. It is one form of dynamic hashing.

Example: Let global depth (gd) = 2. It means the directory contains four entries. Let the local depth (ld) of each bucket = 2. It means each bucket needs two bits to perform search operation. Let each bucket capacity be four. Let us insert 21, 15, 28, 17, 16, 13, 19, 12, 10, 24, 25 and 11.

21 = 101 01	19 = 100 11
15 = 011 11	12 = 011 00
28 = 111 00	10 = 010 10
17 = 100 01	24 = 110 00
16 = 100 00	25 = 111 01
13 = 011 01	11 = 010 11

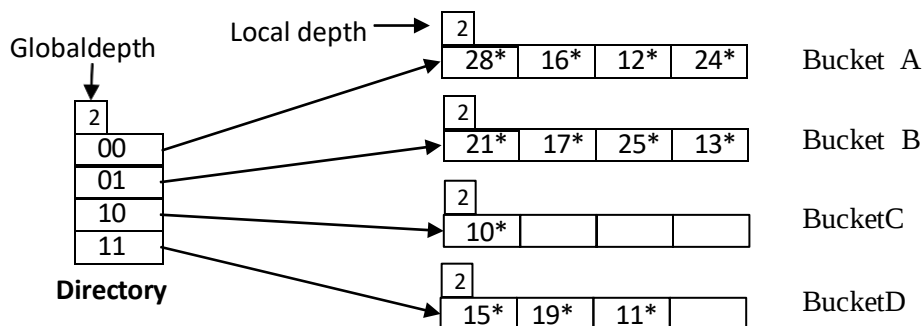


Figure 6.1: Extendable hashing example

Now, let us consider insertion of data entry 20* (binary 101**00**). Looking at the directory element 00, we are led to bucket A, which is already full. So, we have to **split** the bucket by

allocating a new bucket and redistributing the contents (including the new entry to be inserted) across the old bucket and its 'split image (new bucket). To redistribute entries across the old bucket and its split image, we consider the last *three bits*; the last two bits are 00, indicating a data entry that belongs to one of these two buckets, and the third bit discriminates between these buckets. That is if a key's last three bits are 000, then it belongs to bucket A and if the last three bits are 100, then the key belongs to Bucket A2. As we are using three bits for A and A2, so the local depth of these buckets becomes 3. This is illustrated in the below Figure 6.2.

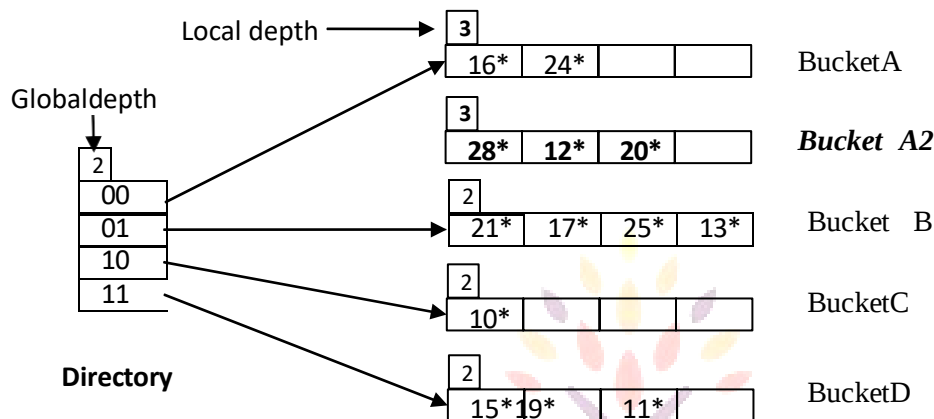


Figure 6.2: After inserting 20 and splitting Bucket A

After split, Bucket A and A2 have local depth greater than global depth ($3 > 2$), so double the directory and use three bits as global depth. As Bucket A and A2 have local depth 3, so they have separate pointers from the directory. But, Buckets B, C and D use only local depth 2, so they have two pointers each. This is shown in the below diagram.

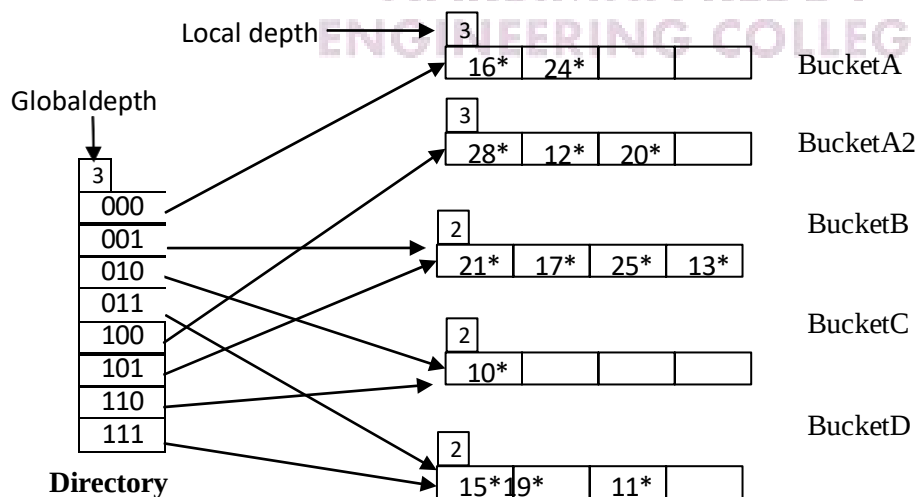


Figure 6.3: After inserting 20 and doubling the directory

An important point that arises is whether splitting a bucket necessitates a directory doubling. Consider our example, as shown in Figure 6.3. If we now insert 9* (01001), it belongs in bucket B; this bucket is already full. We can deal with this situation by splitting the bucket and using directory elements 001 and 101 to point to the bucket and its split image. This is shown in the below diagram. As Bucket B and its split image now have local depth three and it is not greater than global depth. Hence, a bucket split does not necessarily require a directory doubling. However, if either bucket A or A2 grows full and an insert then forces a bucket split, we are forced to double the directory once again.

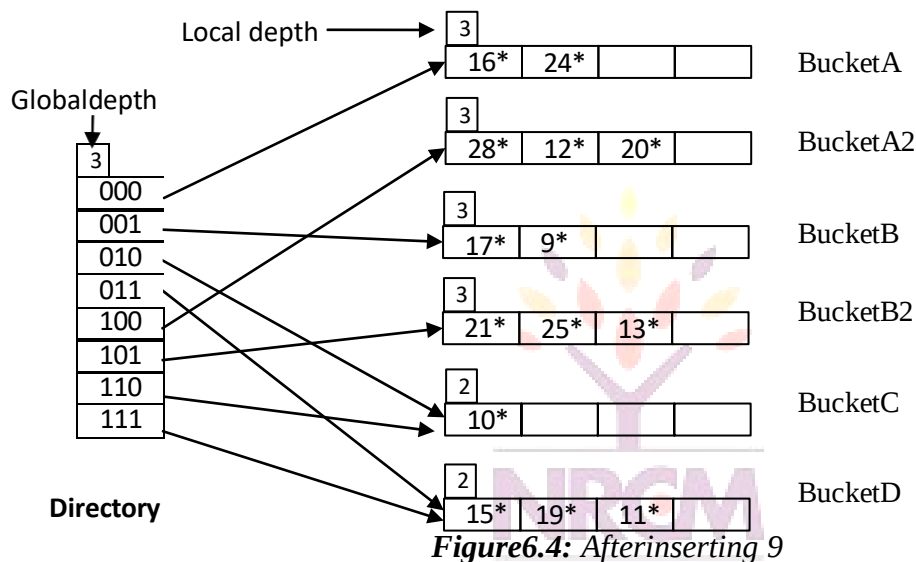


Figure 6.4: After inserting 9

Key Observations:

- A Bucket will have more than one pointers pointing to it if its local depth is less than the global depth.
- When over flow condition occurs in a bucket, all the entries in the bucket are hashed with a new local depth.
- If new Local Depth of the over flowing bucket is equal to the global depth, only then the directories are doubled and the global depth is incremented by 1.
- The size of a bucket cannot be changed after the data insertion process begins.

ii. Linear Hashing

Linear hashing is a dynamic hashing technique that linearly grows or shrinks number of buckets in a hash file without a directory as used in *Extendible Hashing*. It uses a family of hash functions instead of single hash function.

DATABASE MANAGEMENT SYSTEM (25CS305)

This scheme utilizes a *family* of hash functions $h_0, h_1, h_2, \dots$, with the property that each function's range is twice that of its predecessor. That is, if h_i maps a data entry into one of N buckets, h_{i+1} maps a data entry into one of $2N$ buckets. One example of such hash function family can be obtained by:

$$h_{i+1}(\text{key}) = \text{key} \bmod (2^i N)$$

where N is the initial number of buckets and $i = 0, 1, 2, \dots$

Initially it uses N buckets labelled 0 through $N-1$ and an initial hashing function $h_0(\text{key}) = \text{key} \% N$ is used to map any key into one of the N buckets. For each overflow bucket, one of the buckets in serial order will be split and its content is redistributed between it and its split image. That is, for first time overflow in any bucket, bucket 0 will be split, for second time overflow in any bucket; bucket 1 will be split and so on. When number of buckets becomes $2N$, then this marks the end of splitting round 0. Hashing function h_0 is no longer needed as all $2N$ buckets can be addressed by hashing function h_1 . In new round namely splitting-round 1, buckets split once again starts from bucket 0. A new hash function h_2 will be used. This process is repeated when the hash file grows.

Example: Let $N = 4$, so we use 4 buckets and hash function $h_0(\text{key}) = \text{key} \% 4$ is used to map any key into one of the four buckets. Let us initially insert 4, 13, 19, 25, 14, 24, 15, 18, 23, 11, 16, 12 and 10. This is shown in the below figure.

Bucket# $h_1 h_0$				Primary pages				Overflow pages			
0		000	00	4*	24*	16*	12*				
1		001	01	13*	25*						
2		010	10	14*	18*	10*					
3		011	11	19*	15*	23*	11*				

Next, when 27 is inserted, an overflow occurs in bucket 3. So, bucket 0 (first bucket) is split and its content is distributed between bucket 0 and new bucket. This is shown in below figure.

Bucket# $h_1 h_0$				Primary pages				Overflow pages			
0		000	00	24*	16*						
1		001	01	13*	25*						
2		010	10	14*	18*	10*					
3		011	11	19*	15*	23*	11*	27*			
4		100	00	4*	12*						

DATABASE MANAGEMENT SYSTEM (25CS305)

Next, when 30,31 and 34 is inserted, an over flow occurs in bucket 2. So, bucket 2 is split and its content is distributed between bucket 1 and new bucket. This is shown in below figure.

Bucket#	h_1	h_0	Primary pages	Overflow pages
0	000	00	24* 16*	
1	001	01	13*	
2	010	10	14* 18* 10* 30*	34*
3	011	11	19* 15* 23* 11*	27* 31*
4	100	00	4* 12*	
5	101	01	25*	

When 32, 35, 40 and 48 is inserted, an over flow occurs in bucket 0. So, bucket 0 is split and its content is distributed between bucket 2 and new bucket. This is shown in below figure.

Bucket#	h_1	h_0	Primary pages	Overflow pages
0	000	00	24* 16* 32* 40*	48*
1	001	01	13*	
2	010	10	18* 10* 34*	
3	011	11	19* 15* 23* 11*	27* 31* 35*
4	100	00	4* 12*	
5	101	01	25*	
6	110	10	14* 30*	

When 26, 20 and 42 is inserted, an over flow occurs in bucket 0. So, bucket 0 is split and its content is distributed between bucket 3 and new bucket. This is shown in below figure.

Bucket#	h_1	h_0	Primary pages	Overflow pages
0	000	00	24* 16* 32* 40*	48*
1	001	01	13*	
2	010	10	18* 10* 34* 26*	42
3	011	11	19* 11* 27* 35*	
4	100	00	4* 12* 20*	
5	101	01	25*	
6	110	10	14* 30*	
7	111	11	15* 23* 31*	

This marks the end of splitting round. Hashing function h_0 is no longer needed as all $2N$ buckets can be addressed by hashing function h_1 . In new round namely splitting-round 1, bucket split once again starts from bucket 0. A new hash function h_2 will be used. This process is repeated.

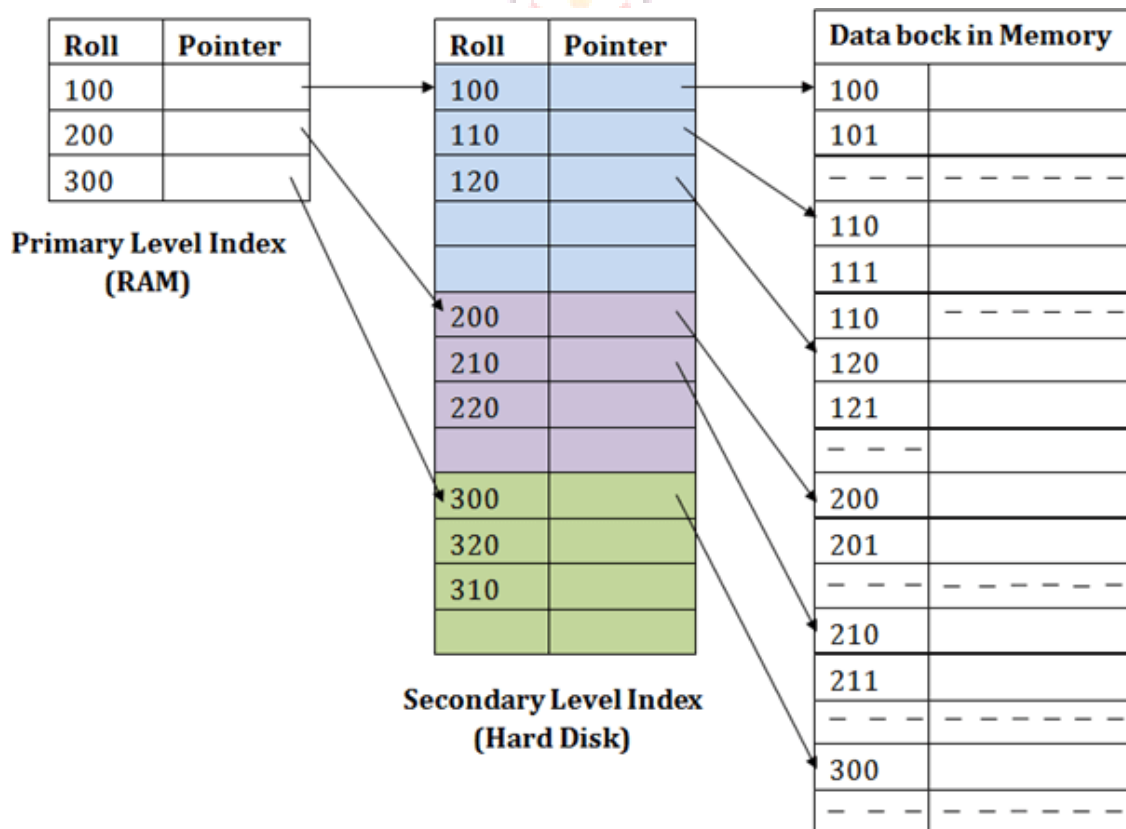


TS	●	→	TS	Hyderabad	KCR
TN	●	→	AP	Amaravathi	Jagan
MH	●	→	TN	Madras	PalaniSwamy
		→	MH	Bombay	Thackray

i. Secondary Index

When the size of the main table grows, then the size of the index table also grows. If the index table size grows then fetching the address itself becomes slower. To overcome this problem, secondary indexing is introduced.

- In secondary indexing, to reduce the size of mapping, another level of indexing is introduced.
- It contains two levels. In the first level, each record in the main table has one entry in the first-level index table.
- The index entries in the first-level index table are divided into different groups. For each group, one index entry is created and added in the second-level index table.

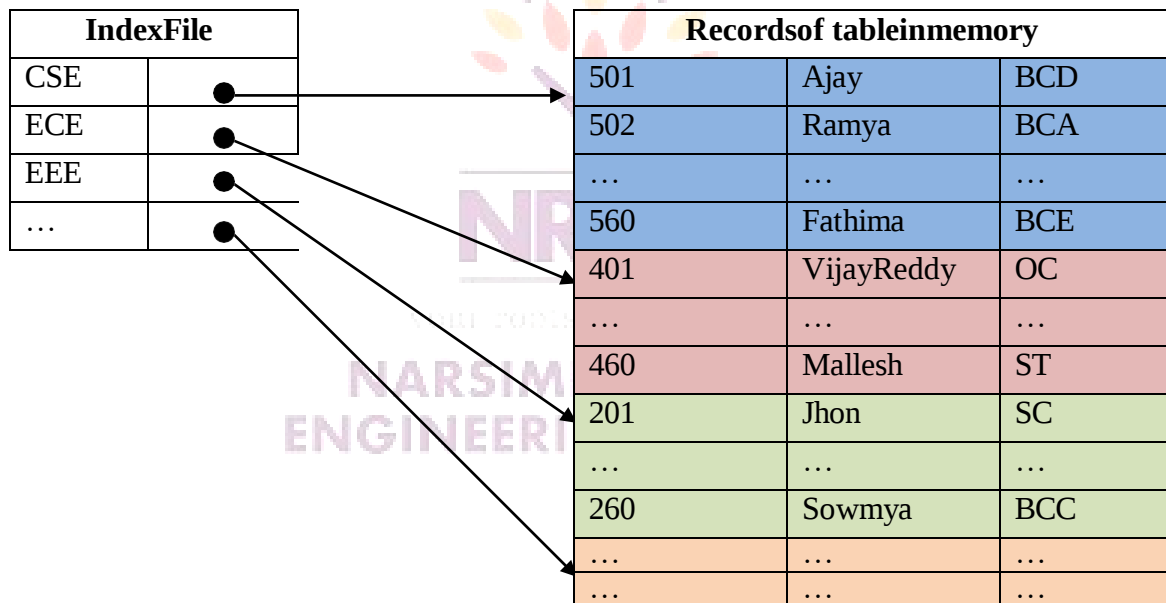


Multi-level Index: When the main table size becomes too large, creating a secondary level index improves the search process. Even if the search process is slow; we can add one more level of indexing and so on. This type of indexing is called multi-level index.

ii. Clustering Index

- Some times the index is created on non-primary key columns which may not be unique for each record.
- In this case, to identify the record faster, we will group two or more columns to get the unique value and create index out of them. This method is called a clustering index.
- The records which have similar characteristics are grouped, and indexes are created for these group.

Example: Consider a college contains many students in each department. All the students belong to the same Dept_ID are grouped together and treated as a single cluster. One index pointers point to the one cluster as a whole. The index pointer points to the first record in each cluster. Here Dept_ID is a non-unique key.



In above diagram we can see that, indexes are created for each department in the index file. In the data block, the students of each department are grouped together to form the cluster. The address in the index file points to the beginning of each cluster.

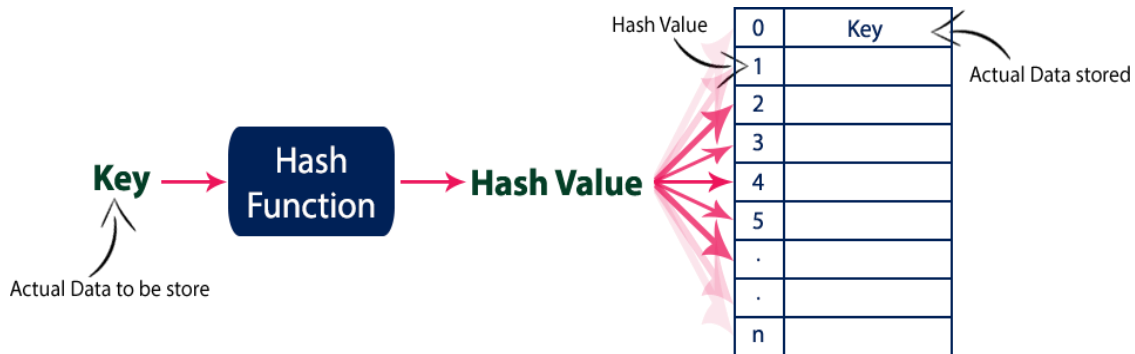
5. HASHBASED INDEXING

Hashing is a technique to directly search the location of desired data on the disk without using index structure. Hash function is a function which takes a piece of data (key) as input and

produces a hash value as output which maps the data to a particular location in the hash table.



The concept of hashing and hash table is shown in the below figure



There are mainly two types of hashing methods:

- i. Static Hashing
- ii. Dynamic Hashing
 - Extended hashing
 - Linear hashing

6. STATIC HASHING

In static hashing, the hash function produces only fixed number of hash values. For example consider the hash function

$$f(x) = x \bmod 7$$

For any value of x , the above function produces one of the hash value from $\{0, 1, 2, 3, 4, 5, 6\}$. It means static hashing maps search-key values to a fixed set of bucket addresses.

Example: Inserting 10, 21, 16 and 12 in hash table.

	Hash Value	Data Record
$f(10) = 10 \bmod 7 = 3$	0	21*
$f(21) = 21 \bmod 7 = 0$	1	
$f(16) = 16 \bmod 7 = 2$	2	16*
$f(12) = 12 \bmod 7 = 5$	3	10*
	4	
	5	12*
	6	

Figure 5.1: Static hashing

Suppose, latter if we want to insert 23, it produce hash value as 2 ($23 \bmod 7 = 2$). But, in the above hash table, the location with hash value 2 is not empty (it contains 16*). So, a collision occurs. To resolve this collision, the following techniques are used.

- Open addressing
- SeparateChainingorClosed addressing

iii. Open Addressing:

Open addressing is a collision resolving technique which stores all the keys inside the hash table. No key is stored outside the hash table. Techniques used for open addressing are:

- Linear Probing
- Quadratic Probing
- Double Hashing

➤ LinearProbing:

In linear probing, when there is a collision, we scan forwards for the next empty slot to fill the key's record. If you reach last slot, then start from beginning.

Example: Consider figure 1. When we try to insert 23, its hash value is 2. But the slot with 2 is not empty. Then move to next slot (hash value 3), even it is also full, then move once again to next slot with hash value 4. As it is empty store 23 there. This is shown in the below diagram.

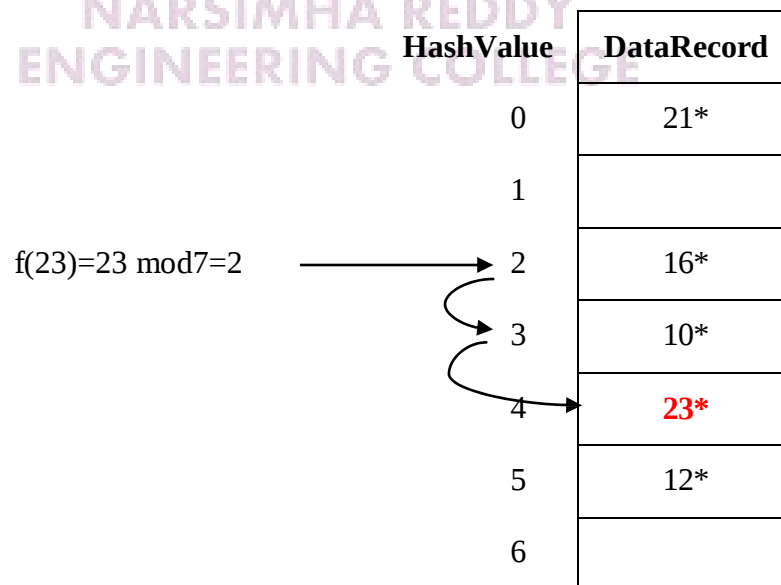


Figure5.2: Linear Probing

➤ Quadratic Probing:

In quadratic probing, when collision occurs, it compute new hash value by taking the original hash value and adding successive values of quadratic polynomial until an open slot is found. If there is a collision, it use the following hash function: $h(x) = (f(x) + i^2) \bmod n$, where $i = 1, 2, 3, 4, \dots$ and $f(x)$ is initial hash value.

Example: Consider figure 1. When we try to insert 23, its hash value is 2. But the slot with hash value 2 is not empty. Then compute new hash value as $(2 + 1^2) \bmod 7 = 3$, even it is also full, and then once again compute new hash value as $(2 + 2^2) \bmod 7 = 6$. As it is empty store 23 there. This is shown in the below diagram.

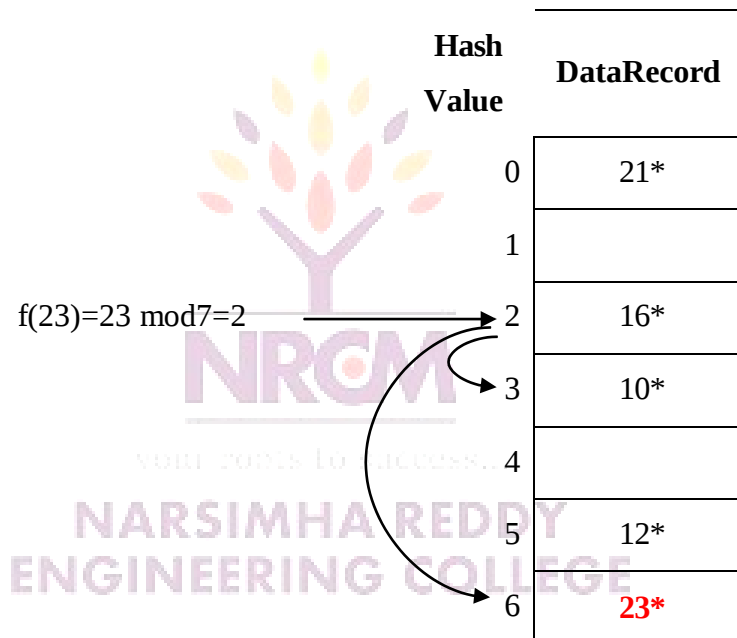


Figure 5.3: Quadratic Probing

➤ Double Hashing

In double hashing, there are two hash functions. The second hash function is used to provide an offset value in case the first function causes a collision. The following function is an example of double hashing: $(\text{firstHash}(\text{key}) + i * \text{secondHash}(\text{key})) \% \text{tableSize}$. Use $i = 1, 2, 3, \dots$

A popular second hash function is : $\text{secondHash}(\text{key}) = \text{PRIME} - (\text{key} \% \text{PRIME})$
 where PRIME is a prime smaller than the TABLE_SIZE.

Example: Consider figure1. When we try to insert 23, its hash value is 2. But the slot with hash value 2 is not empty. Then compute double hashing value as

$$\text{secondHash (key)} = \text{PRIME} - (\text{key} \% \text{PRIME}) \rightarrow \text{second Hash (23)} = 5 - (23 \% 5) = 2$$

$$\text{Double hashing: } (\text{firstHash(key)} + i * \text{secondHash(key)}) \% \text{tableSize} \rightarrow (2 + 1 * 2) \% 7 = 4$$

As the slot with hash value 4 is empty, store 23 there. This is shown in the below diagram.

HashValue	DataRecord
0	21*
1	
2	16*
3	10*
4	23*
5	12*
6	

Figure 5.4: Double Probing

iv. Separate Chaining or Closed addressing:

To handle the collision, This technique creates a linked list to the slot for which collision occurs. The new key is then inserted in the linked list. These linked lists to the slots appear like chains. So, this technique is called as *separate chaining*. It is also called as closed addressing.

Example: Inserting 10, 21, 16, 12, 23, 19, 28, 30 in hash table.

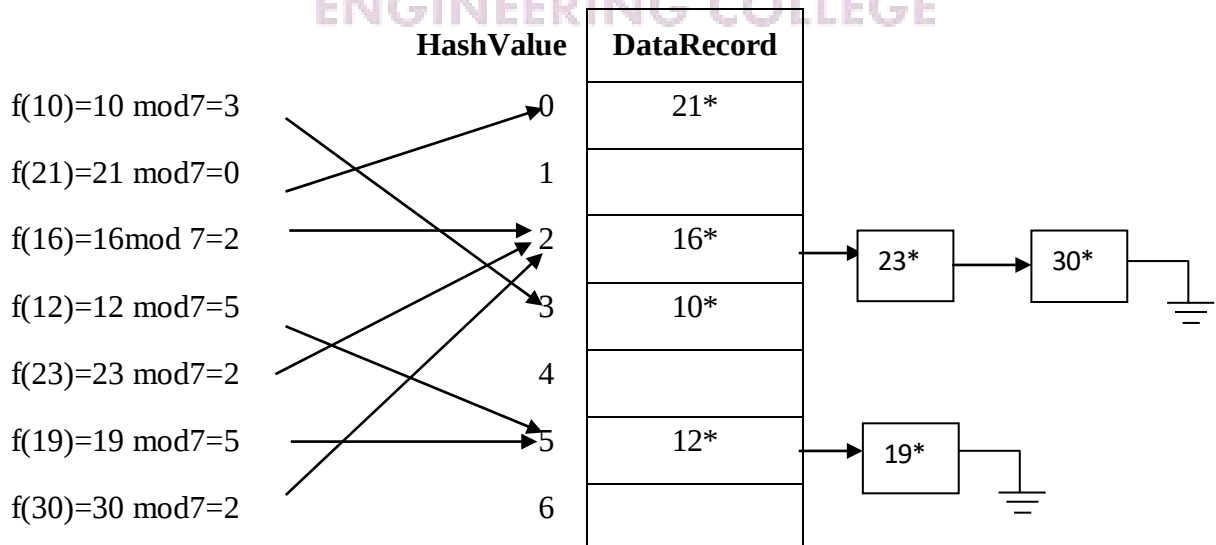


Figure 5.5: Separate chaining example

7. DYNAMIC HASHING

The problem with static hashing is that it does not expand or shrink dynamically as the size of the database grows or shrinks. Dynamic hashing provides a mechanism in which data buckets are added and removed dynamically and on-demand. Dynamic hashing can be implemented using two techniques. They are:

- Extendedhashing
- LinearHashing

iii. Extendable hashing

In extendable hashing, a separate *directory* of pointers to buckets is used. The number bits used in directory is called global depth (gd) and number entries in directory = 2^{gd} . Number of bits used for locating the record in the buckets is called *local depth* (ld) and each bucket can store up to 2^{ld} entries. The hash function uses last few binary bits of the key to find the bucket. If a bucket overflows, it splits, and if local depth greater than global depth, then the table doubles in size. It is one form of dynamic hashing.

Example: Let global depth (gd) = 2. It means the directory contains four entries. Let the local depth (ld) of each bucket = 2. It means each bucket needs two bits to perform search operation. Let each bucket capacity be four. Let us insert 21, 15, 28, 17, 16, 13, 19, 12, 10, 24, 25 and 11.

21 = 101 01	19 = 100 11
15 = 011 11	12 = 011 00
28 = 111 00	10 = 010 10
17 = 100 01	24 = 110 00
16 = 100 00	25 = 111 01
13 = 011 01	11 = 010 11

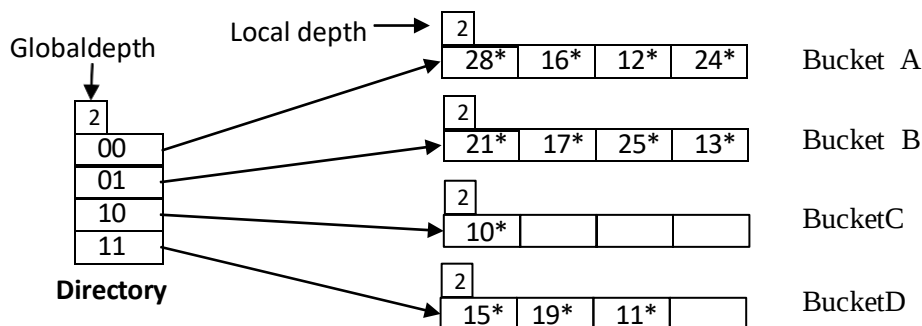


Figure 6.1: Extendable hashing example

Now, let us consider insertion of data entry 20* (binary 101**00**). Looking at directory element 00, we are led to bucket A, which is already full. So, we have to **split** the bucket by

allocating a new bucket and redistributing the contents (including the new entry to be inserted) across the old bucket and its 'split image (new bucket). To redistribute entries across the old bucket and its split image, we consider the last *three bits*; the last two bits are 00, indicating a data entry that belongs to one of these two buckets, and the third bit discriminates between these buckets. That is if a key's last three bits are 000, then it belongs to bucket A and if the last three bits are 100, then the key belongs to Bucket A2. As we are using three bits for A and A2, so the local depth of these buckets becomes 3. This is illustrated in the below Figure 6.2.

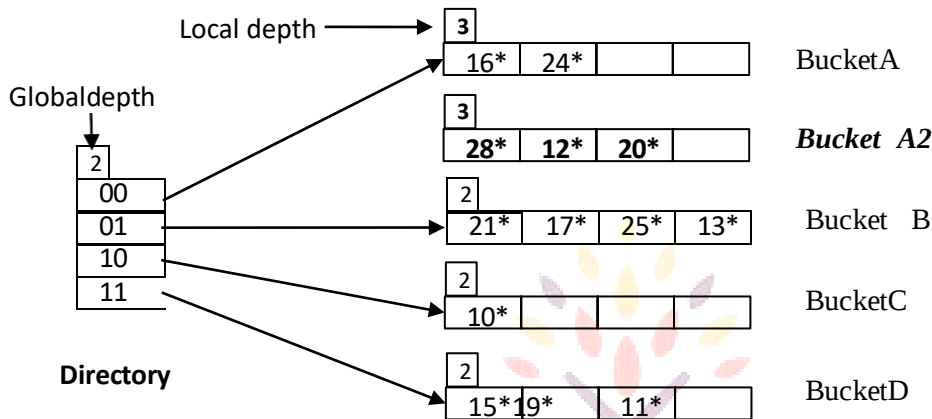


Figure 6.2: After inserting 20 and splitting Bucket A

After split, Bucket A and A2 have local depth greater than global depth ($3 > 2$), so double the directory and use three bits as global depth. As Bucket A and A2 has local depth 3, so they have separate pointers from the directory. But, Buckets B, C and D use only local depth 2, so they have two pointers each. This is shown in the below diagram.

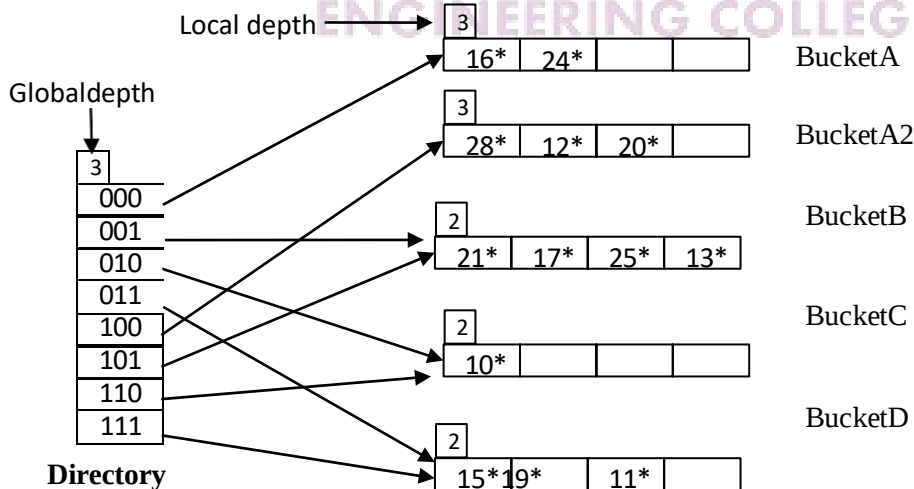


Figure 6.3: After inserting 20 and doubling the directory

An important point that arises is whether splitting a bucket necessitates a directory doubling. Consider our example, as shown in Figure 6.3. If we now insert 9* (01001), it belongs in bucket B; this bucket is already full. We can deal with this situation by splitting the bucket and using directory elements 001 and 101 to point to the bucket and its split image. This is shown in the below diagram. As Bucket B and its split image now have local depth three and it is not greater than global depth. Hence, a bucket split does not necessarily require a directory doubling. However, if either bucket A or A2 grows full and an insert then forces a bucket split, we are forced to double the directory once again.

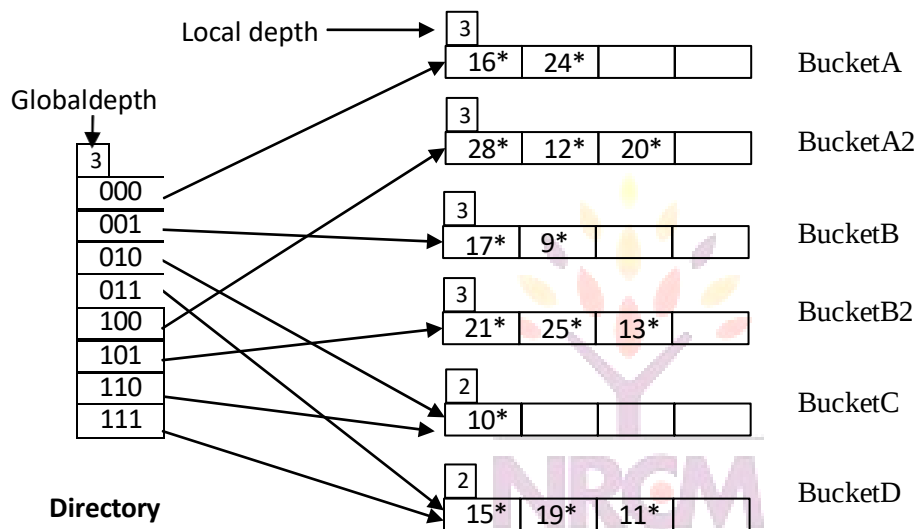


Figure 6.4: After inserting 9

Key Observations:

- A Bucket will have more than one pointers pointing to it if its local depth is less than the global depth.
- When over flow condition occurs in a bucket, all the entries in the bucket are hashed with a new local depth.
- If new Local Depth of the overflowing bucket is equal to the global depth, only then the directories are doubled and the global depth is incremented by 1.
- The size of a bucket cannot be changed after the data insertion process begins.

iv. Linear Hashing

Linear hashing is a dynamic hashing technique that linearly grows or shrinks number of buckets in a hash file without a directory as used in *Extendible Hashing*. It uses a family of hash functions instead of single hash function.

DATABASE MANAGEMENT SYSTEM (25CS305)

This scheme utilizes a *family* of hash functions $h_0, h_1, h_2, \dots$, with the property that each function's range is twice that of its predecessor. That is, if h_i maps a data entry into one of N buckets, h_{i+1} maps a data entry into one of $2N$ buckets. One example of such hash function family can be obtained by:

$$h_{i+1}(\text{key}) = \text{key} \bmod (2^i N)$$

where N is the initial number of buckets and $i = 0, 1, 2, \dots$

Initially it uses N buckets labelled 0 through $N-1$ and an initial hashing function $h_0(\text{key}) = \text{key} \% N$ is used to map any key into one of the N buckets. For each overflow bucket, one of the buckets in serial order will be split and its content is redistributed between it and its split image. That is, for first time overflow in any bucket, bucket 0 will be split, for second time overflow in any bucket; bucket 1 will be split and so on. When number of buckets becomes $2N$, then this marks the end of splitting round 0. Hashing function h_0 is no longer needed as all $2N$ buckets can be addressed by hashing function h_1 . In new round namely splitting-round 1, buckets split once again start from bucket 0. A new hash function h_2 will be used. This process is repeated when the hash file grows.

Example: Let $N = 4$, so we use 4 buckets and hash function $h_0(\text{key}) = \text{key} \% 4$ is used to map any key into one of the four buckets. Let us initially insert 4, 13, 19, 25, 14, 24, 15, 18, 23, 11, 16, 12 and 10. This is shown in the below figure.

Bucket# $h_1 h_0$				Primary pages	Overflow pages
0		000	00	4* 24* 16* 12*	
1		001	01	13* 25*	
2		010	10	14* 18* 10*	
3		011	11	19* 15* 23* 11*	

Next, when 27 is inserted, an overflow occurs in bucket 3. So, bucket 0 (first bucket) is split and its content is distributed between bucket 0 and new bucket. This is shown in below figure.

Bucket# $h_1 h_0$				Primary pages	Overflow pages
0		000	00	24* 16*	
1		001	01	13* 25*	
2		010	10	14* 18* 10*	
3		011	11	19* 15* 23* 11*	27*
4		100	00	4* 12*	

DATABASE MANAGEMENT SYSTEM (25CS305)

Next, when 30, 31 and 34 is inserted, an over flow occurs in bucket 2. So, bucket 1 is split and its content is distributed between bucket 1 and new bucket. This is shown in below figure.

Bucket#	h_1	h_0	Primary pages	Overflow pages
0	000	00	24* 16* 	
1	001	01	13* 	
2	010	10	14* 18* 10* 30*	34*
3	011	11	19* 15* 23* 11*	27* 31*
4	100	00	4* 12* 	
5	101	01	25* 	

When 32, 35, 40 and 48 is inserted, an over flow occurs in bucket 0. So, bucket 2 is split and its content is distributed between bucket 2 and new bucket. This is shown in below figure.

Bucket#	h_1	h_0	Primary pages	Overflow pages
0	000	00	24* 16* 32* 40*	48*
1	001	01	13* 	
2	010	10	18* 10* 34* 	
3	011	11	19* 15* 23* 11*	27* 31* 35*
4	100	00	4* 12* 	
5	101	01	25* 	
6	110	10	14* 30* 	

When 26, 20 and 42 is inserted, an over flow occurs in bucket 0. So, bucket 3 is split and its content is distributed between bucket 3 and new bucket. This is shown in below figure.

Bucket#	h_1	h_0	Primary pages	Overflow pages
0	000	00	24* 16* 32* 40*	48*
1	001	01	13* 	
2	010	10	18* 10* 34* 26*	42
3	011	11	19* 11* 27* 35*	
4	100	00	4* 12* 20* 	
5	101	01	25* 	
6	110	10	14* 30* 	
7	111	11	15* 23* 31* 	

This marks the end of splitting round. Hashing function h_0 is no longer needed as all $2N$ buckets can be addressed by hashing function h_1 . In new round namely splitting-round 1, bucket split once again starts from bucket 0. A new hash function h_2 will be used. This process is repeated.

