

Database Systems II B: DBMS-Implementation

Chapter 13: Query Evaluation

Prof. Dr. Stefan Brass

Martin-Luther-Universität Halle-Wittenberg

Wintersemester 2021/22

<http://www.informatik.uni-halle.de/~brass/dbi21/>

Objectives

After completing this chapter, you should be able to:

- explain what a query evaluation plan (QEP) is.
- explain pipelined evaluation and why sorting needs temporary (disk) space.
- explain different algorithms for implementing joins.

Especially nested loop join and merge join.

- read and explain Oracle QEPs.

If a query performs poorly, you need to be able to understand why.

- develop different query evaluation plans for a given query and assess their merits.

Introduction (2)

- In this chapter, we use a standard example database from Oracle about employees and departments:

```
EMP(EMPNO, ENAME, JOB, SAL, MGR→EMP, DEPTNO→DEPT)
DEPT(DEPTNO, DNAME, LOC)
```

- Consider the following SQL query:

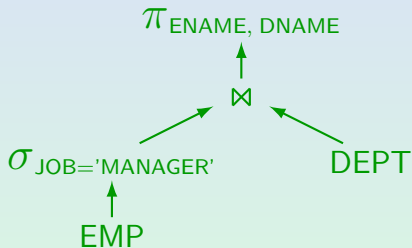
```
SELECT ENAME, DNAME
FROM   EMP, DEPT
WHERE  EMP.DEPTNO = DEPT.DEPTNO
AND    JOB = 'MANAGER'
```

- In relational algebra, this is:

$$\pi_{ENAME, DNAME}(\sigma_{JOB='MANAGER'}(EMP) \bowtie DEPT)$$

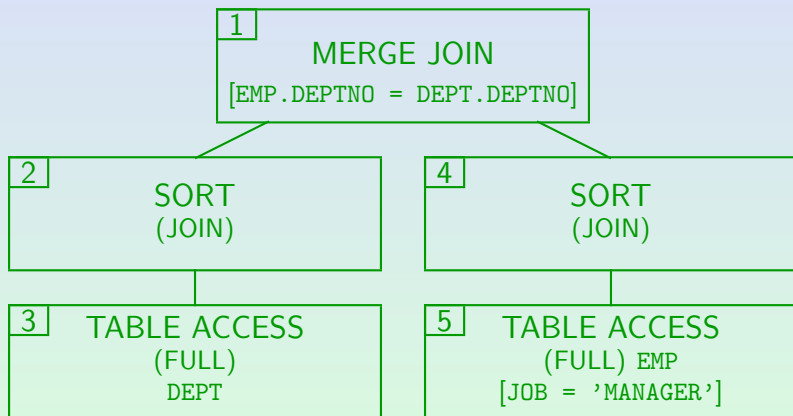
Introduction (3)

- Complex relational algebra expressions are best displayed as “operator trees”:



This shows the flow of data. One can view relations/tuples as being pushed from the base relations in the leaf nodes through the relational algebra operators towards the root, where the final result is computed.

Table 1 Demographic characteristics of study population



(Oracle does not show the small annotations in [...].)

Journal of Management Inquiry 20(6) 798–814



Viewing Oracle QEPs (1)

- First, create a table “**PLAN_TABLE**” in which Oracle will store information about the QEP.

The table must exist under the account of each user who wants to view QEPs. It has prescribed columns, see slide 13-12 for details.

- The simplest way to do this is to execute the script
\$ORACLE_HOME/rdbms/admin/utlxplan.sql

- Then enter the following command in SQL*Plus:
SET AUTOTRACE ON EXPLAIN

Then Oracle will show information about the QEPs for all following queries (not all details, only the structure). If one logs out from SQL*Plus, the AUTOTRACE is forgotten, but the PLAN_TABLE still exists.

Viewing Oracle QEPs (2)

- The output you get from AUTOTRACE is not in graphical form as shown above, but in textual form:

Execution Plan

```
0      SELECT STATEMENT Optimizer=CHOOSE
1      0      MERGE JOIN
2      1      SORT (JOIN)
3      2      TABLE ACCESS (FULL) OF 'DEPT'
4      1      SORT (JOIN)
5      4      TABLE ACCESS (FULL) OF 'EMP'
```

The first number identifies the tree node (shown above in the upper left corner), the second number is the parent node.

Details: Plan Table (1)

EXPLAIN PLAN command:

- An alternative to “**SET AUTOTRACE ON**” is to use **EXPLAIN PLAN FOR** `<SQL QUERY>`

Then Oracle prints only “Explained”. It does not execute the query and does not automatically show the QEP. But information about the QEP is stored in the PLAN_TABLE (can be retrieved with SQL). The rows should normally be deleted before the next EXPLAIN PLAN.

- The PLAN_TABLE can contain rows for several QEPs, then one should use e.g.

```
EXPLAIN PLAN SET STATEMENT_ID = 'MyFirstQuery'
FOR SELECT ... FROM ... WHERE ...
```

Details: Plan Table (2)

- The `PLAN_TABLE` contains one row for each node in the QEP(s) stored in it.

More precisely, not the QEP is stored in it, but only some information about the general QEP structure. Oracle does not show all details of the QEP (e.g. selection conditions).

- Columns of the `PLAN_TABLE`:

- **STATEMENT_ID**: Used to distinguish the rows belonging to execution plans for different queries.

Normally the `PLAN_TABLE` contains only one plan and `STATEMENT_ID` is null. But see `SET STATEMENT_ID` above.

- **TIMESTAMP**: Time when `EXPLAIN PLAN` was issued.

Details: Plan Table (4)

- Columns of the PLAN_TABLE, continued:
 - The following columns describe operations.
 - **OPERATION**: E.g. "TABLE ACCESS", "MERGE JOIN".
 - **OPTIONS**: E.g. "FULL" for operation "TABLE ACCESS".
 - **OBJECT_OWNER**, **OBJECT_NAME**: Identifies the table or index used in the operation.

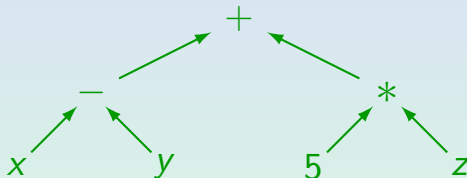
Null for operations which get input only from their children.
 - **OBJECT_INSTANCE**: Position of table in the FROM-list.

E.g. useful if there are two tuple variables over one table.
 - **OBJECT_TYPE**: "UNIQUE"/"NON-UNIQUE" for indexes.

Processing of QEPs (2)

- Such a tree representation is known for arithmetic expressions.

For example: $(x - y) + 5 * z$



- Arithmetic expressions are usually compiled by using registers as temporary storage:

```
R1 := x - y;  
R2 := 5 * z;  
R1 := R1 + R2; // R1 contains now result
```

Processing of QEPs (3)

- It would be possible to
 - compute every operation at once completely,
 - store the result in a temporary relation, and
 - let the parent operation read this relation.
- This corresponds to the compilation of arithmetic expressions with registers as temporary storage for intermediate results:

```
R1 :=  $\sigma_{\text{JOB}='MANAGER'}(\text{EMP})$ ;  
R2 := R1  $\bowtie$  DEPT;  
R3 :=  $\pi_{\text{ENAME}, \text{DNAME}}(\text{R2})$ ;  
print R3;
```

Processing of QEPs (4)

- However, in this way a lot of memory is needed for the intermediate results.
- Sometimes intermediate results are so large that they have to be written to disk and then read again.
- But one can eliminate nearly all temporary storage since most operations work “tuple by tuple”.

Sorting is an exception (see next section).

- In the example, when the join has computed some tuple, one can immediately compute the projection result for that tuple (instead of first storing it).

Processing of QEPs (5)

- Most operations compute tuples only on demand (when the parent node needs them), and only one tuple at a time.

E.g. the join node requests a tuple from the selection node. In order to satisfy the request, it requests a tuple from the relation EMP and checks the condition JOB='MANAGER'. If the condition is satisfied, it returns the tuple and is done. If not, it requests another tuple from the relation EMP.
- Thus, tuples flow immediately from the child to the parent, even before the child has computed the complete result.
- This is called “Pipelined/Lazy Execution”.

Pipelined Evaluation: Interface (1)

Interface of QEP Nodes (Example):

- The interface is very similar to an SQL cursor.

One opens the relation that is the result of this operation, fetches every tuple in a loop, and closes it. (Other names: “scan”, “iterator”).

- In object-oriented terms, there is an abstract class `QEP_Node`, with subclasses for every kind of operator.

E.g. `QEP_Node_Selection` or `QEP_Node_Merge_Join`.

- Constructor: This creates a new QEP node.
The parameters depend on the type of operation.

E.g. a the constructor for `QEP_Node_Selection` needs the child QEP node and the selection condition.

Pipelined Evaluation: Interface (2)

- **open**: Open input.

This method may have parameters (depending on the type of operation).

E.g. the search values for an index scan.

In this way, Information can also flow down in the tree.

- **next**: Advance input to next tuple.

Returns false if end of input.

- `attr(i)`: Value of *i*-th attribute of current tuple.

This returns a pointer to the attribute value. In this way we avoid constructing an entire new tuple for the result.

- **close**: Close input. It may then be opened again.

Pipelined Evaluation: Interface (3)

Less common operations:

- **save/restore**: Remember the current position in the stream of result tuples / switch back to it.

Needed for merge join if duplicate values on both sides ($\times$ for subset).

- **back**: Switch back to previous result tuple.

This operation is inverse to *next*. Needed for zig-zag nested loop join.

- **num_attrs**: Number of attributes in the result.

- **size/cost**: Estimates for number of tuples in the result and the runtime needed for computing them.

This is useful for query optimization.

Example: Selection (1)

- Suppose we want to implement a simple selection of the form $\sigma_{Attr = Val}(Input)$.

In a real system we must be able to pass any condition on tuples (with $\neg$, $\vee$, $\wedge$ and $<$, $>$, like, is null, ...).

- `QEP_Node_Selection(Input, AttrNo, Val)`:
The constructor stores the three parameters in instance variables (attributes) of this object.

- `open()`:
`Input->open(); // Simply pass to child node`
- `close()`:
`Input->close();`

Example: Selection (2)

- `next()`:

```
bool End_of_Input;  
End_of_Input = Input->next();  
while(!End_of_Input  
      && Input->attr(AttrNo) != Val)  
    End_of_Input = Input->next();  
return(End_of_Input);
```
- `attr(i)`:

```
return(Input->attr(i));
```
- `num_attrs()`:

```
return(Input->num_attrs());
```

Contents

- 1 Query Evaluation Plans
- 2 **Sorting**
- 3 Algorithms for Joins
- 4 Operators in Oracle QEPs

Temporary Storage (1)

- Not all operations can compute their results “on demand”.
- E.g. a sort operation needs to see all input tuples before it can return the first result tuple.

Otherwise it is possible that a tuple which is earlier in the sort order is still to come.

- Thus, a sort operation needs temporary space for storing all input tuples.

Temporary Storage (2)

- Of course, the sort operation has the same external interface as all other QEP nodes.

open, next, close, ...

- However:
 - During the open, it will already read and sort all its input tuples (i.e. the real work is done here).
 - Then later requests for the next result tuple will be answered from the intermediate storage.

Temporary Storage (3)

- Sometimes it would also be good to materialize other intermediate results which have to be read more than once (e.g. in a nested loop join).

Some systems have a special operator for doing this (“Bucket”). But Oracle seems to use intermediate space only for sorting.

- In Oracle, the maximal size of temporary storage that a single sort operation can request in memory is set by the initialization parameter `SORT_AREA_SIZE`.
- If the space needed for sorting is larger, Oracle will use temporary segments on disk.

Temporary Storage (4)

- The current value of this parameter is shown with:

```
SELECT VALUE
FROM   V$PARAMETER
WHERE  NAME = 'sort_area_size';
```

On our UNIX systems, the default is 65536 Bytes.

- The parameter can be changed with

```
ALTER SESSION SET SORT_AREA_SIZE = 131072;
```

The memory is taken from the Program Global Area (PGA), i.e. inside the dedicated server process, not from the SGA. However, in the multithreaded server configuration, it is taken from the SGA.

Temporary Storage (5)

- After the sort is done, the sorted rows must be temporarily stored until they are fetched.
- `SORT_AREA_RETAINED_SIZE` controls how much memory can be used for this purpose.

By default, this parameter is the same as `SORT_AREA_SIZE`. But if memory is scarce, it should be used for running sorts rather than afterwards when the rows only wait to be fetched.

- There are more initialization parameters controlling the sorting.

```
SELECT NAME, VALUE, DESCRIPTION FROM V$PARAMETER
WHERE NAME LIKE '%sort%'
```

Temporary Storage (6)

- Temporary segments can be allocated in any tablespace, but it is better to use a special “temporary tablespace”.

The storage parameters for the temporary segments are inherited from the tablespace in which they are allocated. `INITIAL` should be a multiple of the `SORT_AREA_SIZE` plus one block for the segment header.

- The tablespace used for temporary segments can be defined separately for each user.

See `CREATE USER` statement. It can be changed with `ALTER USER`.

- Information about temporary segments is available in `V$SORT_SEGMENT` and `V$SORT_USAGE`.

Performance Statistics (1)

- How many sorts in the current session were done in memory? How many on disk? And how many rows were sorted?

```
SELECT X.VALUE, Y.NAME
FROM   V$SESSTAT X, V$STATNAME Y, V$SESSION Z
WHERE  X.STATISTIC# = Y.STATISTIC#
AND    Y.NAME LIKE '%sort%'
AND    X.SID = Z.SID AND Z.USERNAME = USER
```

There is also a table V\$SYSSTAT which contains accumulated counts since the DBMS was last started. These statistics are also contained in the report produced by utlstat.sql/utlstat.sql (see above).

Performance Statistics (2)

- After **SET AUTOTRACE ON**, SQL*Plus prints not only the QEP for every query, but also performance statistics (including information about sorts).

SET AUTOTRACE ON STATISTICS prints only the statistics.

- The role **PLUSTRACE** gives access to some dynamic performance views. It must be granted to all users who should be able to use this feature.

It contains access to `sys.v_$sesstat`, `sys.v_$statname`, `sys.v_$session`. To declare this role, the DBA (user SYS) must execute the script `plustrce.sql`. It is located in `$ORACLE_HOME/sqlplus/admin`.

Sort Algorithm (1)

- Sorting is needed quite often, and it is a relatively expensive operation.
- Thus, many thoughts were put into developing an efficient sort algorithm, and new improvements are still proposed in the literature.
- Sorting with external memory is usually based on the merge sort algorithm, which you should know from your data structures course.

Sort Algorithm (2)

- Mergesort is based on the notion of “runs”, which are already sorted sequences of elements.
- E.g. when you want to sort n elements, you start with n runs of length 1.
- Then you always merge two such sorted sequences (“runs”) of length l to one sorted sequence of length $2 * l$.

Sort Algorithm (3)

- The merging can be done in linear time: You look at the first element of both runs, take the smaller one and put it into the output. Repeat this until both runs are empty.
- Since the size of the runs doubles every time, you need a logarithmic number of iterations until you have only one run which contains all elements. → Complexity $O(n * \log(n))$.

You can implement it with four files: Two for the input runs and two for the output runs. Output runs are written to the two files in alternating fashion so that they contain the same number of runs.

Example (Basic Mergesort)

- Input (16 runs of length 1):

12	5	9	20	16	18	3	7	17	10	2	25	13	15	6	8
----	---	---	----	----	----	---	---	----	----	---	----	----	----	---	---

- After first step (8 runs of length 2):

5	12	9	20	16	18	3	7	10	17	2	25	13	15	6	8
---	----	---	----	----	----	---	---	----	----	---	----	----	----	---	---

- Second and third step:

5	9	12	20	3	7	16	18	2	10	17	25	6	8	13	15
---	---	----	----	---	---	----	----	---	----	----	----	---	---	----	----

3	5	7	9	12	16	18	20	2	6	8	10	13	15	17	25
---	---	---	---	----	----	----	----	---	---	---	----	----	----	----	----

- After fourth step (1 runs of length 16: final result):

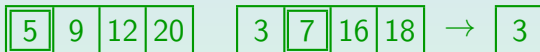
2	3	5	6	7	8	9	10	12	13	15	16	17	18	20	25
---	---	---	---	---	---	---	----	----	----	----	----	----	----	----	----

Example (Merging of Runs)

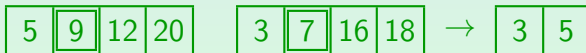
- One first compares the first elements of both runs:



- 3 is smaller, so it is written to the output and the current position in the second file is moved forward:



- Now 5 is smaller, and written to the output:



And so on (exercise). When the end of file is reached on one side, the rest of the other side is written to the output.

Sort Algorithm Optimizations

- There are many optimizations to Mergesort, e.g.:
 - One tries to produce large initial runs by sorting chunks of the given elements in the available main memory. The longer the initial runs, the less iterations are needed later.

Once a block of such an initial run was written to disk, one can reuse the memory page for more input elements. New elements which happen to be greater than the greatest element already written to the output can still become part of the current run.

- If one has k buffer frames available during the merge phase, one merges $k - 1$ runs instead of only 2 runs.

Contents

- 1 Query Evaluation Plans
- 2 Sorting
- 3 Algorithms for Joins
- 4 Operators in Oracle QEPs

Nested Loop Join (1)

- The nested loop join
 - looks at all combinations of tuples from both relations,

This is done (in the most simple version) by an outer loop over the tuples of the left relation, in which an inner loop on the right relation is nested, see the algorithm on the next slide.
 - evaluates the join condition, and
 - returns those combinations for which the condition is true.
- $R \bowtie_{A_i=B_j} S$ is evaluated similarly to $\sigma_{A_i=B_j}(R \times S)$ but without materializing the intermediate result of $\times$.

Our pipelined evaluation anyway would not materialize the result, but we nevertheless save many function calls.

Nested Loop Join (2)

- Without the pipelined evaluation, the algorithm for $R \bowtie_{A_i=B_j} S$ looks as follows:

```

(1)  foreach tuple  $t = (d_1, \dots, d_n)$  in  $R$  do
(2)      foreach tuple  $u = (e_1, \dots, e_m)$  in  $S$  do
(3)          if  $d_i = e_j$  then
(4)              output  $t \circ u = (d_1, \dots, d_n,$ 
(5)                   $e_1, \dots, e_m)$ ;
(6)          fi;
(7)      od;
(8)  od;
```

- Thus the name “nested loop”.

Nested Loop Join (3)

- If both relations have approximately n tuples each, n^2 tuple combinations are checked.
- Thus, the nested loop join needs quadratic time, i.e. its complexity is $O(n^2)$.
- The merge join (see below) is asymptotically faster: It has complexity $O(n * \log(n))$.
- However, the nested loop join works for arbitrary join conditions, not only equality conditions.

The merge join and other specialized join methods work only with equality conditions like $A = B$.

Nested Loop Join: Optimizations (1)

- If the left (outer) relation has ℓ tuples, the right (inner) relation will be read ℓ times.
- This can be improved by reading as many as possible tuples from the left relation into memory, and then doing the join with all these tuples during one pass through the inner relation.

If n tuples from the left relation fit into the available memory, the right one will be read only ℓ/n times.

- It is better to use the smaller relation as left (outer) relation.

It must be read once in addition to the $\ell * r$ or $r * \ell$ tuples read in the inner loop.

Nested Loop Join: Optimizations (2)

- If one of the two relations is small enough to fit into main memory, each relation must be read only once from disk.

Since the smaller relation is used as outer relation and the outer relation gets all available buffer space, this happens automatically in the above scheme. Note that still $\ell * r$ tuple combinations have to be considered, only the number of tuples read from disk shrinks to $\ell + r$.

- Even if only one buffer frame can be used for the inner relation, it should be read forward and backward in alternating fashion (zig-zag). This gives one block access less in every pass (starting from the second).

Nested Loop Join: Optimizations (3)

- If the left (outer) input relation is sorted, the unoptimized version would protect the sorting. The optimized version destroys it.

If the other (unsorted) relation fits into main memory, one can use it as inner relation to protect the sorting on the outer one.

- Since there are multiple passes through the inner input relation, it might be useful to store it explicitly (not using pipelined evaluation) as compact as possible.
- This discussion should demonstrate that even for a relatively simple operation like the nested loop join, there are many optimization tricks.

Merge Join (1)

- The merge join works very similar to merge sort.
- Both input relations must be sorted on the join attribute.
- Then the algorithm does a parallel pass on both relations:
 - It advances always the scan with the smaller value in the join attribute.

That value cannot have a join partner on the other side, since all following values there will be even bigger than the current one.
 - In this way it finds all matches (equal values).

Merge Join (2)

$R \bowtie_{A_i=B_j} S$:

```

(1)  open( $R$ ); open( $S$ );
(2)  read  $t = (d_1, \dots, d_n)$  from  $R$ ;
(3)  read  $u = (e_1, \dots, e_m)$  from  $S$ ;
(4)  while not eof( $R$ ) and not eof( $S$ ) do
(5)      if  $d_i < e_j$  then
(6)          read  $t = (d_1, \dots, d_n)$  from  $R$ ;
(7)      else if  $d_i > e_j$  then
(8)          read  $u = (e_1, \dots, e_m)$  from  $S$ ;
(9)      else /*  $d_i = e_j$  */
(10)         output  $t \circ u = (d_1, \dots, d_n, e_1, \dots, e_m)$ ;
(11)         read  $u = (e_1, \dots, e_m)$  from  $S$ ;

```

This program code assumes that A_i is a key in R . Therefore, after a match is found, the other side S is advanced for a possible further match.

Example (Merge Join)

Selection on EMP		
<u>EMPNO</u>	ENAME	DEPTNO
7782	CLARK	2
7839	KING	2
7934	MILLER	2
7369	SMITH	3
7876	JONES	3
7788	SCOTT	3
7566	ADAMS	6
7499	ALLEN	7
7654	MARTIN	7

DEPT	
<u>DEPTNO</u>	DNAME
1	ACCOUNTING
2	RESEARCH
3	SALES
4	OPERATIONS
7	SHIPPING

ADAMS violates the foreign key, but makes the example more interesting.

Merge Join (4)

- The time needed for the join itself is linear in the size of the two relations.

We assume again that the join attribute on one side is a key of that relation, so there are no duplicate values on that side. If duplicate values were allowed on both sides, the extreme case (a single value repeated n times) would always lead to quadratic complexity: This would simply be a kartesian product.

- If we have to sort them, the total complexity is $O(n * \log(n))$.

In comparison, the runtime (CPU time) of the nested loop join is always quadratic in the sizes of the input relations. The number of block accesses is only quadratic if neither one fits into memory. However, the merge join works only for equality conditions.

Index Join

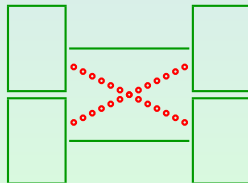
- Suppose we have to compute $R \bowtie_{A=B} S$ and that there is an index on $S(B)$.
- Then we can loop over all tuples in R and locate the corresponding tuples from S via the index.
- Since every access to S via the index potentially needs one or more block accesses, this is only useful if R contains only relatively few tuples (less tuples than S has blocks).

Otherwise the merge sort is better.

The index is also useful if S is small and will be completely buffered, but then there probably should be no index.

Hash Join (1)

- The idea of the hash join is to partition both relations into small pieces by applying a hash function to the join attribute.
- Possible matches can only occur between tuples with the same hash value. Only such tuple combinations must be tried, not all tuple combinations.



Hash Join (2)

- The partitioning is done such that the smaller parts fit into main memory.
- If needed, the partitioning step is iterated.
- Then a hash table is built in memory for each such partition and an index join is done with the corresponding partition of the other table.
- The result is the union of the joins of the pairs of partitions with the same hash value.

Hash Join (3)

Example: $\text{hash}(\text{row}) = \begin{cases} 1 & \text{if DEPTNO odd} \\ 2 & \text{otherwise} \end{cases}$

Selection on EMP		
<u>EMPNO</u>	ENAME	DEPTNO
7369	SMITH	3
7499	ALLEN	7
7654	MARTIN	7
7788	SCOTT	3
7782	CLARK	2
7839	KING	2
7566	ADAMS	6
7934	MILLER	2

DEPT	
<u>DEPTNO</u>	DNAME
1	ACCOUNTING
3	SALES
7	SHIPPING
2	RESEARCH
4	OPERATIONS

Contents

- 1 Query Evaluation Plans
- 2 Sorting
- 3 Algorithms for Joins
- 4 Operators in Oracle QEPs**

Optimizer Modes

- Oracle has four possible optimizer modes (see next chapter):
 - **RULE**: Older optimizer which does not depend on statistics about table sizes (rule-based optimizer).
 - **FIRST_ROWS**: Newer cost-based optimizer with the goal to produce the first row fast (best response time).
 - **ALL_ROWS**: Newer cost-based optimizer with the goal to produce the last row fast (best throughput).
 - **CHOOSE**: Use the cost-based approach at least one table used in the statement was analyzed.
- **ALTER SESSION SET OPTIMIZER_MODE = ...;**
- The QEPs shown below were generated with the rule-based optimizer.

Oracle QEPs for Simple Queries (1)

- Simple Query = Join, Selection, and Projection.
- Above we have looked at query evaluation plans for:

```

SELECT ENAME, DNAME
FROM    EMP, DEPT
WHERE   EMP.DEPTNO = DEPT.DEPTNO
  
```

- It should now be understandable what the merge join means and why it requires that both inputs are sorted.
- The other access plan (page 4-6) is really an index join but uses Oracle's NESTED LOOPS operator.
- It seems that Oracle tries to evaluate conditions as far down in the tree as possible (which is good).

Oracle QEPs for Simple Queries (2)

- So Oracle's NESTED LOOPS operator does not itself evaluate the join condition.
- E.g. when the current employee in the outer loop works in department 20, the condition DEPTNO=20 is passed to the right child when it is opened.
- Also the TABLE ACCESS operator in Oracle can contain a selection which is not explicitly shown. E.g. the QEP for the following query consists only of the TABLE ACCESS node:

```
SELECT EMPNO  
FROM EMP  
WHERE ENAME LIKE 'F%'
```

Oracle QEPs for Simple Queries (3)

- Suppose we change the join query by adding a condition on EMP:

```

SELECT ENAME, DNAME
FROM   EMP, DEPT
WHERE  EMP.DEPTNO = DEPT.DEPTNO
AND    EMP.SAL > 5000

```

- Here Oracle would show the same access plan as without the condition `EMP.SAL > 5000`.
- But of course, this condition is evaluated in the `TABLE ACCESS` operator before the sorting for the merge join.
- In general, one tries to evaluate selections as early as possible to make the intermediate results smaller.

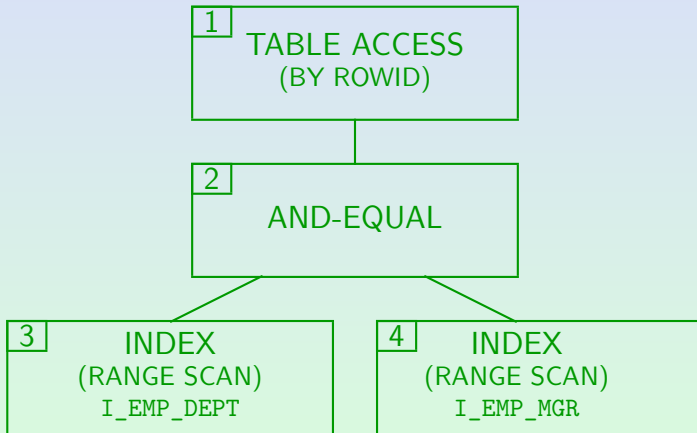
Oracle QEPs for Simple Queries (4)

- Consider the following query:

```
SELECT *  
FROM   EMP  
WHERE  DEPTNO = 20 AND MGR = 7566
```

- If there are indexes on both attributes, Oracle produces an evaluation plan which intersects the ROWIDs returned from both indexes (see next page).
- These access plans are computed by the rule-based optimizer.
- If you analyze the tables EMP and DEPT, the cost-based optimizer becomes applicable. It creates a full table scan in this case since the table EMP fits into one block.

Oracle QEPs for Simple Queries (5)



Oracle QEPs for Simple Queries (6)

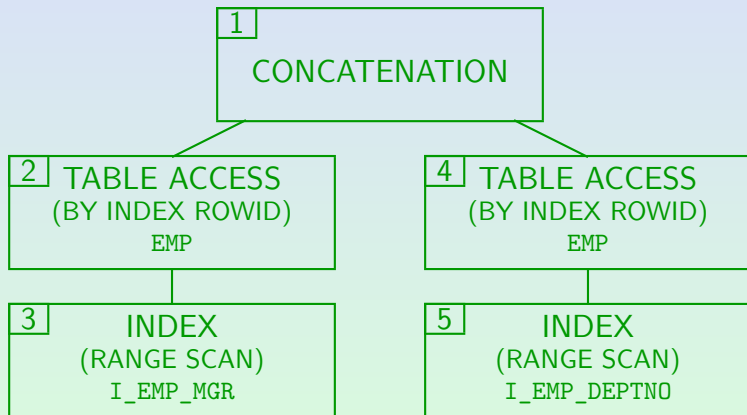
- Consider the following query:

```
SELECT * FROM EMP
WHERE DEPTNO = 20 OR MGR = 7566
```

- If there are indexes on both attributes, the rule-based optimizer produces an evaluation plan which uses both indexes and concatenates the results (see next page).
- Tuples which satisfy both conditions are printed only once (as it should be) although the cheap concatenation does no duplicate elimination. Probably it is evaluated this way:

```
SELECT * FROM EMP WHERE MGR = 7566
UNION ALL
SELECT * FROM EMP WHERE DEPTNO = 20
AND NOT (MGR = 7566)
```

Oracle QEPs for Simple Queries (7)



Oracle QEPs for Simple Queries (8)

- Oracle can produce index-only QEPs:

```
SELECT EMPNO
FROM   EMP
WHERE  MGR = 7566
```

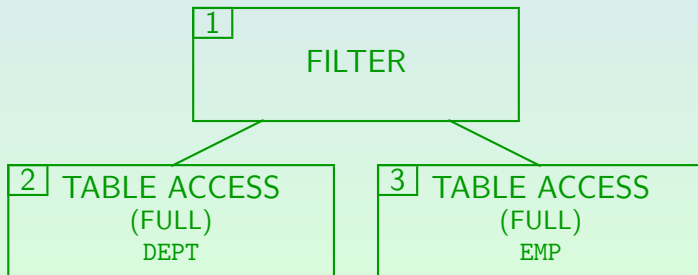
- If there is an index on EMP(MGR, EMPNO), Oracle will produce this plan:



- However, if there is an index on EMP(MGR), Oracle will prefer this, even though afterwards it has to access the table.

Oracle QEPs for NOT EXISTS (1)

```
SELECT *  
FROM   DEPT D  
WHERE  NOT EXISTS(SELECT * FROM EMP E  
                  WHERE E.DEPTNO = D.DEPTNO)
```



Oracle QEPs for NOT EXISTS (2)

- **FILTER** is an operation that accepts a set of rows, eliminates some of them, and returns the rest.

See Oracle8 Tuning, page 23-8.

- So FILTER is something like a selection, but Oracle never shows proper selections explicitly, they are done as a by-product in the other nodes.
- FILTER is only needed to evaluate complex conditions, especially with subqueries (e.g. NOT EXISTS).
- FILTER may have any number of child nodes (at least two). The left child of FILTER computes the rows which are to be filtered. All other childs correspond to subqueries.

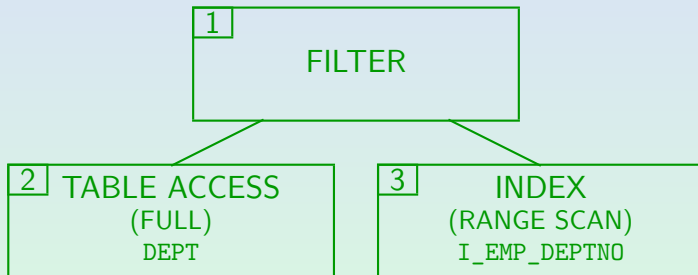
Oracle QEPs for NOT EXISTS (3)

- Consider again the query for departments without employees:

```
SELECT *
FROM   DEPT D
WHERE  NOT EXISTS(SELECT * FROM EMP E
                  WHERE E.DEPTNO = D.DEPTNO)
```

- If there is an index I_EMP_DEPTNO on EMP(DEPTNO), Oracle produces an access plan which does not access the relation EMP, but only the index (see next page).

1000

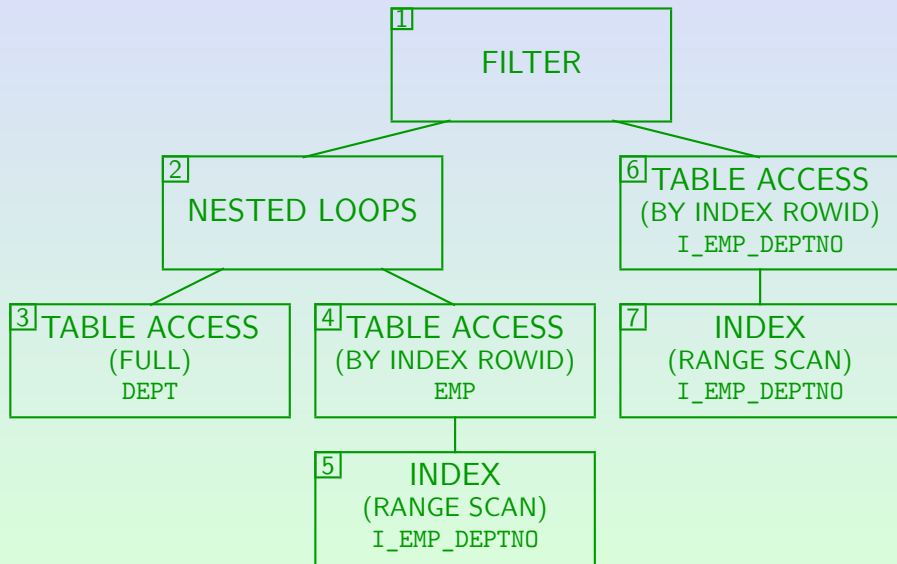


Oracle QEPs for NOT EXISTS (5)

- The set of rows which have to be filtered can also be computed by a join.
- E.g. consider the following query which computes employees in the research department which do not have a supervisor in that department (the QEP is shown on the next page):

```
SELECT E.ENAME
FROM   EMP E, DEPT D
WHERE  E.DEPTNO = D.DEPTNO
AND    D.DNAME = 'RESEARCH'
AND    NOT EXISTS (SELECT * FROM EMP X
                   WHERE  X.EMPNO = E.MGR
                   AND     X.DEPTNO = D.DEPTNO)
```

Oracle QEPs for NOT EXISTS (6)



Oracle QEPs for NOT EXISTS (7)

- Oracle also has Anti-Join Operations which can be used for NOT EXISTS and NOT IN subqueries.
- An antijoin works like a join, but returns those tuples from the first relation which have no join-partner in the second relation.
- Oracle has anti-join versions of the merge join and the hash join. It seems that FILTER is something like an anti-join version of a nested loop join.
- You can request to use a specific kind of anti-join with the initialization parameter ALWAYS_ANTI_JOIN.

But this parameter is not modifiable while the system runs.

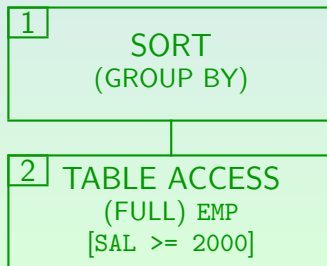
Oracle QEPs for Aggregations (1)

- This query returns the number of employees per department:

```

SELECT  DEPTNO, COUNT(*)
FROM    EMP
WHERE   SAL >= 2000
GROUP BY DEPTNO
  
```

- The input relation (of course, only tuples satisfying the WHERE-condition) is sorted by the GROUP BY attribute(s):



Oracle QEPs for Aggregations (2)

- A query without aggregation will use the operator
SORT (AGGREGATE):

```
SELECT COUNT(*)  
FROM EMP  
WHERE SAL >= 2000
```

- This operator does not sort, but computes the single tuple which is the result of the aggregation.
- A query with DISTINCT uses the operator
SORT (UNIQUE):

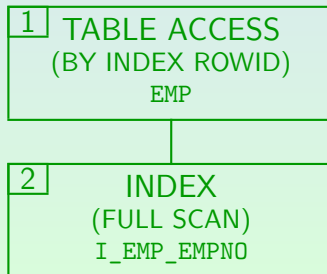
```
SELECT DISTINCT DEPTNO  
FROM EMP
```

Oracle QEPs for Sorting (1)

- Consider a simple query with ORDER BY:

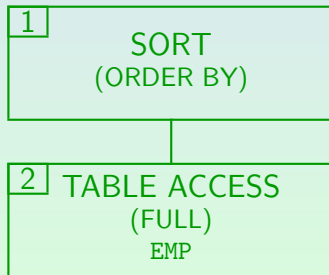
```
SELECT      *  
FROM        EMP  
ORDER BY EMPNO
```

- Oracle can use an index for producing a sorted sequence:



Oracle QEPs for Sorting (2)

- But this is only possible if the sort attribute is NOT NULL. Otherwise, not all tuples are represented in the index.
- E.g. ORDER BY MGR needs a full table scan, even if there is an index on EMP(MGR):



Operations in Oracle QEPs (1)

Access Methods (Tables):

- Access methods are used for getting data from the disk (as opposed to operations which combine data, e.g. joins).
- **TABLE ACCESS (FULL)**: Retrieval of all rows from a table by reading all blocks below the high water mark.
- **TABLE ACCESS (BY ROWID)**: Retrieval of rows from a table given their ROWIDs.
- **TABLE ACCESS (CLUSTER)**: Retrieval of rows from a table stored in an index cluster, given the value of the cluster key.
- **TABLE ACCESS (HASH)**: Retrieval of rows from a table stored in a hash cluster, given the value of the cluster key.

Operations in Oracle QEPs (2)

Access Methods (Indexes):

- **INDEX (UNIQUE SCAN)**: Retrieval of a single ROWID from an index for a key, given a value for the key attribute.
- **INDEX (RANGE SCAN)**: Retrieval of attribute values and ROWIDs from an index, given an interval for the values of the indexed attributes (or a single attribute value).
- **INDEX (RANGE SCAN DESCENDING)**: As before, but attribute values are produced in descending order.

I have not seen any Oracle QEP using this operation.

- **INDEX (FULL SCAN)**: Retrieval of all values for the indexed column in sorted order (plus the corresponding ROWIDs).

Operations in Oracle QEPs (4)

Join Methods:

- **NESTED LOOPS**: Nested Loops or Index Join.
- **NESTED LOOPS (OUTER)**: The same method, but used for an outer join (protecting tuples from one or both tables).
- **MERGE JOIN**: Merge Join (needs sorted inputs).
- **MERGE JOIN (OUTER)**: Outer join based on merge method.
- **MERGE JOIN (SEMI)**: The special case of a semi-join (only checking the existence of a join partner).
- **MERGE JOIN (ANTI)**: The merge method applied for finding tuples without a join partner (for NOT EXISTS, NOT IN).

Operations in Oracle QEPs (5)

Join Methods, Continued:

- **HASH**: A hash join (partitioning both tables).
- **HASH (SEMI)**: The special case of a semi-join.
- **HASH (ANTI)**: The hash method applied for an anti-join.

Sort Operations:

- **SORT (JOIN)**: Sorting before a merge join.
- **SORT (UNIQUE)**: Sorting and duplicate elimination.
- **SORT (GROUP BY)**: Sorting, grouping, aggregation.
- **SORT (AGGREGATE)**: Aggregation to a single row (no sort).
- **SORT (ORDER BY)**: Sorting because of an ORDER BY clause.

Operations in Oracle QEPs (6)

Set Operations:

- **AND-EQUAL**: Intersection of ROWIDs returned from different indexes.
- **CONCATENATION**: Union without duplicate elimination.
- **UNION-ALL**: The same (?).

The documentation also mentions a UNION with duplicate elimination, but I have seen no QEP which uses it. For queries containing UNION, Oracle does a UNION-ALL followed by SORT (UNIQUE).

- **MINUS**: Set difference (needs sorted inputs).
- **INTERSECTION**: Set intersection (needs sorted inputs).

Operations in Oracle QEPs (8)

Other Operations:

- **FILTER**: Evaluate complex conditions (e.g. NOT EXISTS). This is similar to a NESTED LOOPS anti-join.
- **CONNECT BY**: For evaluating queries on tree-structured data with the CONNECT BY clause (not in the SQL-92 standard).
- **PARTITION (CONCATENATED)**: For partitioned tables.
- **INLIST ITERATOR**: Iterates over the operation below it.
- **PROJECTION**: Projection (appears seldom explicitly).
- **VIEW**: Used to mark subqueries which correspond to a view.
- **FOR UPDATE**: Locks the rows flowing through this operator.

QEP Structural Rules

- **TABLE ACCESS (FULL)** and **INDEX** cannot have child nodes.
- **TABLE ACCESS (BY INDEX ROWID)** must have exactly one child node which produces ROWIDs.
- **NESTED LOOPS/MERGE JOIN** must have exactly two children.
- If the query contains n tuple variables, you need $n - 1$ join operators. In the 1999 exam, some students used more.
- **NESTED LOOPS** passes the information about the current tuple from the left child down to its right child etc. Only there an index on the join attribute can be used.

Some students used the index on the join attribute in the left child, or in the right child of a merge join.

References

- Elmasri/Navathe: Fundamentals of Database Systems, 3rd Ed., Chap. 18: “Query Processing and Optimization”
- Silberschatz/Korth/Sudarshan: Database System Concepts, 3rd Ed., Chap. 12: “Query Processing”
- Ramakrishnan/Gehrke: Database Management Systems, 2nd Ed., Mc-Graw Hill, 2000, Chap. 11: “External Sorting”, Chap. 12: “Evaluation of Relational Operators”,
- Kemper/Eickler: Datenbanksysteme (in German), Chap. 8, Oldenbourg, 1997.
- Härder/Rahm: Datenbanksysteme — Konzepte und Techniken der Implementierung (in German), Springer, 1999.
- Garcia-Molina/Ullman/Widom: Database System Implementation. Prentice Hall, 1999, ISBN 0130402648, 672 pages.
- Oracle 8i Concepts, Release 2 (8.1.6), Oracle Corporation, 1999, Part No. A76965-01. Chapter 21: “The Optimizer”.
- Oracle 8i Designing and Tuning for Performance, Release 2 (8.1.6), Oracle Corporation, 1999, Part No. A76992-01.
- Lipeck: Skript zur Vorlesung Datenbanksysteme (in German), Univ. Hannover, 1996.