



Query Optimization in the Presence of Limited Access Patterns

Daniela Florescu, Alon Levy, Ioana Manolescu, Dan Suciu

► To cite this version:

Daniela Florescu, Alon Levy, Ioana Manolescu, Dan Suciu. Query Optimization in the Presence of Limited Access Patterns. [Research Report] RR-3639, INRIA. 1999. inria-00073035

HAL Id: inria-00073035

<https://inria.hal.science/inria-00073035>

Submitted on 24 May 2006

HAL is a multi-disciplinary open access archive for the deposit and dissemination of scientific research documents, whether they are published or not. The documents may come from teaching and research institutions in France or abroad, or from public or private research centers.

L'archive ouverte pluridisciplinaire **HAL**, est destinée au dépôt et à la diffusion de documents scientifiques de niveau recherche, publiés ou non, émanant des établissements d'enseignement et de recherche français ou étrangers, des laboratoires publics ou privés.

A



INSTITUT NATIONAL DE RECHERCHE EN INFORMATIQUE ET EN AUTOMATIQUE

Query Optimization in the Presence of Limited Access Patterns

Daniela Florescu, Alon Levy, Ioana Manolescu, Dan Suciu

No 3639

Mars 1999

_____ THÈME 3A _____

A large blue rectangle occupies the lower half of the page. Overlaid on the left side of this rectangle is a large, light gray stylized letter 'R'. To the right of the 'R', the words 'Rapport de recherche' are written in a white serif font. A horizontal gray brushstroke underline is positioned below the text.

Rapport
de recherche

Query Optimization in the Presence of Limited Access Patterns

Daniela Florescu*, Alon Levy†, Ioana Manolescu‡, Dan Suciu§

Thème 3A — Interaction homme-machine,
images, données, connaissances
Projet Rodin

Rapport de recherche n° 3639 — Mars 1999 — 27 pages

Abstract: We consider the problem of query optimization in the presence of limitations on access patterns to the data (i.e., when one must provide values for one of the attributes of a relation in order to obtain tuples). We show that in the presence of limited access patterns we must search a space of *annotated query plans*, where the annotations describe the inputs that must be given to the plan. We describe a theoretical and experimental analysis of the resulting search space. Based on the conclusions of this analysis, we describe a novel query optimization algorithm that is designed to perform well under the different conditions that may arise. The algorithm searches the set of annotated query plans, pruning invalid and non-viable plans as early as possible in the search space. The algorithm also uses a best-first search strategy in order to produce a first complete plan early in the search. We describe experiments to illustrate the performance of our algorithm.

Key-words: query optimization, heterogeneous databases, semistructured data, access patterns

(Résumé : *tsvp*)

* Daniela.Florescu@inria.fr

† alon@cs.washington.edu

‡ ioana.Manolescu@inria.fr

§ sucIU@research.att.com

Optimisation de Requetes en Presence des Restrictions d'Acces

Résumé : Nous considérons le problème de l'optimisation des requetes dans le cas des limitations d'accès aux données (par exemple, des cas où on est obligé de fournir des valeurs pour certains attributs d'une relation pour pouvoir obtenir des enregistrements de cette relation). Nous montrons que dans ces cas, l'optimiseur doit chercher le meilleur plan dans un espace de *plans décorés*, où la décoration représente les valeurs des attributs qui doivent etre fournies. Ce travail comporte une analyse théorique et experimentale de l'espace de recherche résultant. En nous appuyant sur les conclusions de cette analyse, nous décrivons un nouvel algorithme d'optimisation de requetes, qui est conçu dans le but d'avoir des performances acceptables dans des situations possibles très differentes. Cet algorithme parcourt l'espace des plans décorés, en éliminant les plans qui ne sont pas exécutable ou invalides aussi tot que possible. Cet algorithme utilise aussi une recherche en profondeur, pour poduire un premier plan complet au plus vite. Nous décrivons aussi des experiments qui montrent les performances de cet algorithme.

Mots-clé : Optimisation de requetes, bases de données hétérogènes, données semi-structurées, restrictions d'accès

1 Introduction

The goal of a query optimizer of a database system is to translate a declarative query expressed on a logical schema into an imperative query execution plan that accesses the physical storage of the data, and applies a sequence of relational operators. In building query execution plans, traditional relational query optimizers try to find the most efficient method for accessing the necessary data. When possible, a query optimizer will use auxiliary data structures such as an index on a file in order to efficiently retrieve a certain set of tuples in a relation. However, when such structures do not exist or are not useful for the given query, the alternative of scanning the entire relation always exists. The existence of the fall back option to perform a complete scan is an important assumption in traditional query optimization.

Several recent query processing applications have the common characteristic that it is not always possible to perform complete scans on the data. Instead, the query optimization problem is complicated by the fact that there are only limited access patterns to the data. One such application is optimization in the presence of foreign and table functions [11, 13, 1]. In most cases, such functions require a set of given inputs in order to return a set of tuples. Two of the application areas that are motivating our work are query processing for data integration and processing of graph-based data.

- **Data integration:** a data integration system needs to access autonomous remote sources in order to answer a query. Even if the data integration system is able to model the contents of the remote sources as relations, the sources may provide only limited access patterns to the data that they serve. This may happen for two main reasons (1) the underlying data may actually be stored in a structured file or legacy system hence the interface to the data is naturally limited, and (2) even if the data is stored in a traditional database system, the source may provide only limited access capabilities for reasons of security or performance. A common example of an access limitation is a web form interface that requires that some of the fields in the form be filled in order to obtain answers.
- **Graph based data models:** several recent applications have found data models based on graphs to be most natural (e.g., biological data, semistructured data, data models for the web, XML [2, 3, 4]). In these applications the database is modeled as labeled directed graph, where the nodes correspond to objects and the labels on the edges correspond to attributes names. The graph is used as a *logical* model on which queries are posed. However, there is no unique best method for storing such graphs, and the storage is usually tuned to the particular application. The particular way of storing the graph poses natural limitations on the access patterns to the data. For example, it is often possible to allow following edges only in the forward direction, and the system may not allow to scan the entire set of objects in the graph. We note that such limitations apply also in object-oriented databases.

In both cases above, the fundamental reason for the limitation on available access patterns is the existence of a *mismatch* between the logical model of the data (which is used for posing queries) and the actual physical storage of the data (used for evaluating queries). In the presence of this mismatch, the query optimizer requires an explicit description of the available access patterns to the data. A wide class of the access pattern limitations to the data can be modeled using *binding patterns*. A binding pattern specifies which attributes of a relation *must* be given values when accessing a set of tuples. For example, the binding pattern $R(A^b, B^f)$ specifies that the only way of retrieving tuples of $R(A, B)$ is by providing values for the attribute A . In fact, binding patterns can be viewed as a method for abstracting the storage of the data for the query optimizer.

This paper considers the problem of query optimization in the presence of limitations on access patterns, described by binding patterns. Specifically, given a set of binding patterns describing the only ways of accessing tuples in a set of relations, and given a select-project-join query over these relations, our task is to efficiently find an optimal query execution plan for the query, if a plan exists. Naturally, the query execution plan must adhere to the binding pattern limitations on the relations.

Before we begin our discussion, we note several closely related problems that are not treated in this paper. As we noted earlier, binding patterns are a method for providing the query optimizer explicit descriptions of the data storage. GMAPs [5, 22] are also a method for describing storage structures. The focus of GMAPs is on describing groupings of the data in a data structure rather than describing limited access patterns to the data. In Section 8 we discuss how these two disjoint classes of storage descriptions can be combined.

We noted the applicability of optimization in the presence of binding patterns to the problem of data integration. It is important to emphasize that binding pattern limitations are only one of several issues that arise in the data integration context. One issue that is not treated in this paper is that of semantic reformulation of user queries [6, 7, 27]. Moreover, the problem of considering binding patterns is one aspect of the problem of considering *source capabilities*. The other aspect of this problem is the case in which sources can perform *local* computations [28, 29]. We purposely isolated the problem of binding patterns from the others to understand it better. In Section 8 we also discuss the problems involved in extending other query optimization paradigms (e.g., transformational, randomized) to accommodate binding-pattern limitations.

1.1 Motivating Example

We begin by illustrating the problem with an example. The example is drawn from an actual application of integration of scientific data. The goal of the application, conducted by the Marine Institute of Crete, is to study the levels of water pollution in the Mediterranean Sea. The application includes two sources of data each containing the results of sets of experiments. The first source stores the results of some experiments concerning water circulation, and the second source reports the results of experiments concerning the level of pollution in the water. Integrating the data from the two sources enables the scientists to predict water pollution levels for a wide range of times and locations, by combining the results of the pollution experiments and those of the water circulation. The data resulting from the experiments is stored in the two sources with the following schemas (note that each experiment is defined by a unique key):

```
Source 1:
  Experi1(key, date, depth)
  Location1(key, location)
  Result1(key, circulation)
Source 2:
  Experi2(key, date, depth)
  Location2(key, location)
  Result2(key, emission)
```

In both of the sources the data is stored in a proprietary data store rather than a relational database. Accordingly, the possible operations on the data are limited. In source 1, it is possible to ask for the keys of all the experiments that have been done at a certain date (i.e., to select on a given date) or at a certain depth. Given an experiment key, it is possible to retrieve its location from relation `Location1`, which is a complex value encoding a geographical coordinates of a rectangle, or to retrieve the result of the experiment (from relation `Result1`) which is a complex picture (picture + list of vectors) describing the speed and direction of the water circulation. In this source it is not possible to perform a selection on a location or on the circulation.

The situation is similar for Source 2. Given a date or a depth we can retrieve the key of the corresponding experiments. Using the key, we can retrieve the location and the pollution emission, which is an complex image. Selections on the location or the emission are not possible.

In both sources, the location represents the geographical coordinates of a rectangle. However, the sharing of the sea surface in rectangles is not done in the same fashion across the two sources, i.e. the sources refer to different sets of rectangles. In order to facilitate the integration, the scientists use a third data source that answers queries about such rectangles. Given two input locations the source returns a number between 0 and 1 describing the similarity between the sources (0 represents disjointness and 1 represents that they're identical).

```
Source 3:
  Coincides(location, location, similarity)
```

The typical operation that scientists need to perform on this data is to answer queries of the form: “retrieve the water circulation and the pollution emission on the 1/10/98 for locations matching with high degree of similarity (=0.9). The query can be written as the following conjunctive query:

```
Query(w1,w2):-Experi1(x1,y1,z1),Location1(x1,t1),Result1(x1,w1),
               Experi2(x2,y1,z1),Location2(x2,t2),Result2(x2,w2),Coincides(t1,t2,s)
               y1='1/10/98', s=0.9.
```

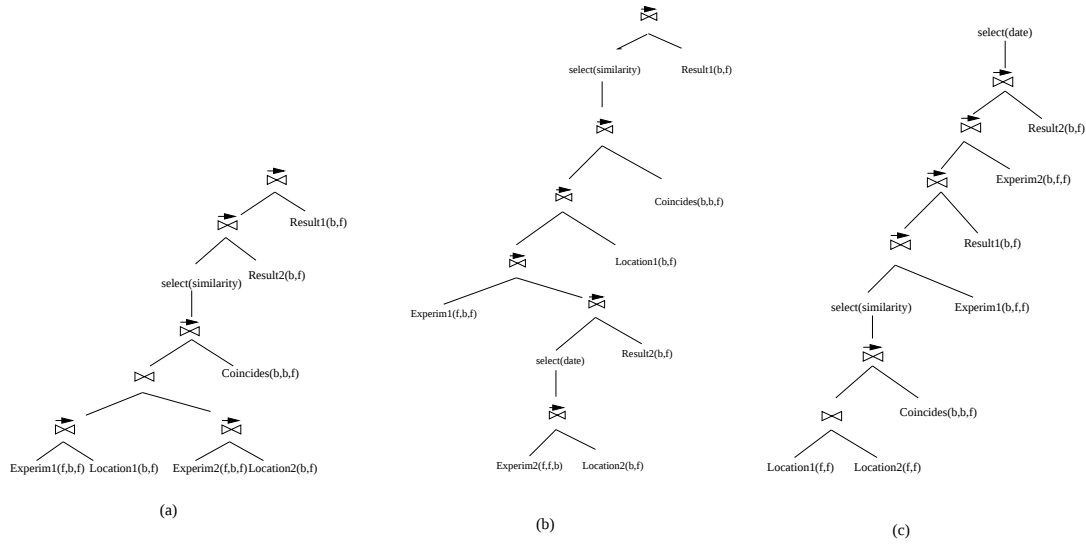


Figure 1: Three different execution plans for the example. (a) and (b) are valid plans, while (c) is not.

Our goal is to find the optimal way to evaluate this query, given the access limitations imposed by the data sources. Figure 1 describes two possible ways to evaluate this query, which are valid according to the limitations and one relational query execution plan which is *not* valid given the access limitations.

The evaluation strategy followed by the plan 1(a) is the following. First, use the selection condition on date on both sources in order to retrieve the keys of the experiments performed on this date. Then, in each source, the keys are used to obtain the corresponding experiment locations. This operation is performed by a dependent join. The join (on depth) of those two temporary relations is sent to Source 3 which calculates the similarity degree for each pair of locations and selects those that satisfy the similarity predicate. In the last step of the evaluation, the keys of the selected experiments are used again to retrieve the desired images (water circulation and pollution emission). The dependent join could be implemented by either of the algorithms described in [10]: one-tuple-at-a-time, once-per-distinct-value, by foreign-sort-merge, etc.

A different plan to evaluate this query, depicted in figure 1(b), is to start by retrieving from Source 1 the keys of the experiments performed on the “1/10/98”, together with their corresponding depth. For each tuple in the result, the values of the depth and date could be used for the following complex computation (which is executed by the engine of the data integration system): using the depth, we can retrieve from the relation `Experim2` the keys of the experiments performed in source 2 on this depth, then the result is filtered using the desired date; the key of the resulting experiments are sent to source 2 in order to retrieve their respective location, and then they are sent once more in order to retrieve the corresponding image. Finally, the location of the experiments of source 1 is retrieved, source 3 is tested for the similarity and, in the last step, the image from source 1 is retrieved. If thousands of experiments are performed per day and per depth, it is reasonable to assume that we cache the result of the inner query execution tree, such that we do not needlessly repeat the same computation for the same values.

Finally, we remark that the plan depicted in figure 1(c) *cannot* be executed given the limitations to the data sources that are given. The reason is that it is impossible to retrieve all the tuples of the relation `Location` in either of the two sources.

1.2 Our Solution

The solution we propose in this paper is based on extending System-R cost-based optimization to incorporate limitations on binding patterns. The key idea underlying our solution is that the optimizer searches through the space of *annotated query plans*, where the annotation of a subplan describes which variables of the query must be given as input to the subplan. We study the effect

of adding annotations on the size of the resulting search space, and describe an efficient algorithm for searching the space.

The idea of adding binding patterns as annotations to subqueries is not new. Such annotations were used in magic-set transformations [8, 9] and for exploring sideways information passing strategies. The focus of this paper is on incorporating such annotations into cost-based optimization and studying the effects of such annotations on the size of the space of query execution plans. Furthermore, System-R also annotates query execution plans with the set of constants necessary for executing the plan. However, the set of necessary constants is always a subset of those given in the query, while our annotations may include variable values passed from other subplans. As we show in our discussion, considering a larger set of annotations may lead to finding better plans.

To summarize, this paper makes the following contributions.

- We show how the presence of binding pattern limitations affects several fundamental properties of the search space, such as the need to consider different binding-pattern annotations on query execution plans, the need to explore the space of bushy trees as opposed to left-linear trees, and the specialized handling of placing selections.

To address these issues, we show how to extend the traditional cost-based query optimization to search through space of *valid* annotated query execution plans. In particular, a query execution plan is annotated with a set of bindings that it receives as inputs in addition to the classical annotations that include subgoals that the plan covers and whether or not the plan produces an interesting order.

- We provide an analytical and empirical study of the effect of adding annotations on the size of the search space. The study considers different shapes of queries, bushy vs. left-linear trees, plans with or without cartesian products, and different numbers of binding patterns associated with each database relation. While the study shows that in some important cases the number of valid query execution plans is actually considerably smaller than the corresponding case without annotations, there are still important cases in which search spaces grows significantly compared to traditional System-R optimization.
- We describe a query optimization algorithm that is designed to perform efficiently under the properties exposed by our analysis. First, the algorithm considers only valid execution plans. The algorithm also prunes early on plans that are not *viable*, i.e., cannot be part of any valid execution plan. Second, the algorithm uses a best-first search strategy in order to produce the first complete query execution plan relatively fast. In contrast, System-R bottom-up optimization only produces a complete query towards the end of the optimization process. Such behavior would not be acceptable in cases in which the search space is significantly larger than in the traditional case. Finally, the algorithm uses a novel method to combine the join enumeration and selection placement.
- We describe an implementation of our optimization algorithm and results of a set of experiments to evaluate its performance. The experiments consider the performance of our algorithm under different conditions (varying query size, shape and number of available binding patterns). The experiments show how our algorithm obtains the first query execution plans much faster than a pure dynamic programming approach. Furthermore, when we consider the time to perform exhaustive search of the space of plans, we show that the extra cost associated with employing a best-first search algorithm compared to dynamic programming only causes linear slowdown w.r.t. the size of the search space. Hence, we argue that a best-first search strategy is more appropriate for the optimization problem we consider.

Section 2 formally defines our problem. Section 3 describes the effects of binding patterns on various properties of the search space, and Section 4 describes an analysis of the size of the resulting search space. Section 5 describes our query optimization algorithm. Section 6 describes our implementation and Section 7 describes the results of our experiments. We end with a discussion and comparison to related work.

2 Problem Definition

In this section we formally define the problem setting we consider in this paper.

Queries: We consider the class of select-project-join queries, also known as *conjunctive queries*. We use the following notation of *conjunctive queries*. A query q is denoted by:

$$q(\bar{X}) : -e_1(\bar{X}_1), \dots, e_n(\bar{X}_n), \mathcal{C}.$$

The predicates $e_1, \dots, e_n$ denote database relations, and $\bar{X}, \bar{X}_1, \dots, \bar{X}_n$ are tuples of variables. The atoms $e_1(\bar{X}_1), \dots, e_n(\bar{X}_n)$ are the conjuncts (or subgoals) of the query, which together with $\mathcal{C}$ form the query's body. The atom $q(\bar{X})$ is the *head* of the query, and the variables in $\bar{X}$ are those that are selected in the result. We require the query to be safe, i.e., any variable that appears in the head must also appear in at least one of the $\bar{X}_i$'s. Note that in this notation, joins predicates are expressed by multiple occurrences of the same variable in different subgoals.

$\mathcal{C}$ is a set of atoms of the form $X_i = c_i$, where X_i appears in $\bar{X}_1 \cup \dots \cup \bar{X}_n$, and c_i is a constant. The set of variables appearing in $\mathcal{C}$ are called the bound variables in the query, denoted by $\text{bound}(q)$.

Data access descriptions: With each database relation we associate a set of binding patterns, describing the possible access patterns to the tuples in the relation. Formally, a binding pattern for a relation R is a mapping from the arguments of R to the alphabet $\{b, f\}$. We depict binding patterns by superscripts on the attribute names of a relation. The meaning of a binding pattern bp for a relation R is that the attributes of R that bp maps to b must be given values when accessing the tuples of R . For example, the binding pattern $R(A^b, B^f)$ for the relation $R(A, B)$ specifies an access pattern where the values of A must be given in order to obtain tuples of R . The traditional flat storage of a relation corresponds to the case where all attributes are mapped to f . A relation may have multiple binding patterns describing the different possible ways to obtain tuples from the relation.

Given a set of binding patterns for a relation R , the query execution engine may choose any of the allowable patterns to access the data.

Example 2.1 The binding patterns for the Example in Section 1.1 are the following.

Source 1:

```

  Experm1(keyf, dateb, depthf)
  Experm1(keyf, datef, depthb)
  Experm1(keyf, dateb, depthb)
  Location1(keyb, locationf)
  Result1(keyb, circulationf)

```

Source 2:

```

  Experm2(keyf, dateb, depthf)
  Experm2(keyf, datef, depthb)
  Experm2(keyf, dateb, depthb)
  Location2(keyb, locationf)
  Result2(keyb, emissionf)

```

Source 3:

```

  Coincides(locationb, locationb, similarityf)

```

Each binding pattern is also labeled with (1) the cost of accessing it once, and (2) the cardinality of the expected output. In our example, if the binding patterns $\text{Experm1}(\text{key}^f, \text{date}^b, \text{depth}^f)$ is labeled with the cost n and the cardinality m , this means that the cost of retrieving the set of tuples in the relation Experm1 given a constant date d is n , and that the number of experiments performed per day is on average m .

Query execution plans: a query execution plan for a query q is a tree whose leaves are labeled with relations in the query and whose internal nodes are algebraic operators. We refer to the leaves of a query execution plan as *atomic plans*. In our discussion we consider plans with join and selection operators. In this paper we consider only selections on the simple form $X_i = c_i$. To simplify our discussion we do not consider plans with projections, and assume they are introduced at a later stage.

We distinguish two kinds of join operators: regular joins and dependent joins. Both types of joins are binary operators and apply recursively on subtrees corresponding to query execution sub-plans. In the case of a regular join, the two input query execution sub-plans can be executed independently of each other, resulting in two tables that can be joined using any of the traditional join algorithms (e.g., hash-join, sort-merge join). In the second type of join, the right input subtree cannot be executed independently, because it requires bindings that are obtained from the result of the left subtree.

Among the algorithms developed for the join operator, only the nest-loop join is applicable to dependent joins. The efficient implementation of dependent joins are considered in [10] in the context

of optimization for foreign functions, as well as in the context of evaluation of path expressions in object-oriented databases. Furthermore, several caching techniques to optimize the implementation of dependent joins are discussed in the context of the Montage system [13]. Finally, we note that that dependent joins have received several other names in the literature (e.g., functional join, implicit join, filter join, theta semi-join, bind join).

As mentioned earlier, our algorithm is going to search the space of *annotated* query execution plans. The annotation describes which variables in the query *must* be given as input to the plan. Formally, each node n in a query execution plan tree is labeled by a pair $(conj(n), adorn(n))$, where $conj(n)$ is the set of conjuncts of q that is covered by n , and the adornment $adorn(n)$ describes which variables of q must be given, in order for the subtree n to be executable.

An adornment is a mapping from the variables in q to the set $\{b, f, _ \}$. The meaning of the adornment is the following: (a) if a variable is mapped to b , then it must be given a value in the start of the execution of the sub-plan, (b) if a variable is mapped to $_$, then it does not appear in the sub-plan corresponding to n and (c) if a variable X is mapped to f , then by executing that sub-plan, we obtain values for X . In our examples $adorn(n)$ is shown as a subscript of $conj(n)$, and shows exactly the set of variables mapped to b . For example, $[R_1(x_0, x_1), R_2(x_1, x_2)]_{\{x_0\}}$ denotes the equivalence class for where the adornment maps x_0 to b , x_1, x_2 to f , and the rest of the variables in the query to $_$.

Example 2.2 The following algebraic expression represents a fragment of the query execution plan depicted in figure 1(a):

$$(Experim1(x_1^f, y_1^b, z_1^f) \bowtie_{x_1} Location1(x_1^b, t_1^f)) \bowtie_{y_1, z_1} (Experim2(x_2^f, y_1^b, z_1^f) \bowtie_{x_2} Location2(x_2^b, t_2^f))$$

This partial plan is annotated by the set of conjuncts and the set of variables:

$$[Experim1(x_1, y_1, z_1), Location1(x_1, t_1), Experim2(x_2, y_1, z_1), Location2(x_2, t_2)]_{\{y_1\}}$$

Clearly, annotations on the nodes in a query execution plan are not arbitrary. Given a set of data access descriptions to the database relations and a query execution plan, the annotations in the plan must satisfy the following conditions. If n is a leaf accessing a relation R , then n should specify which of the available access patterns to R should be used. The adornment of n is correct if it is obtained by extending one of the binding patterns associated with R by mapping each variable not appearing in n to $_$.¹ If n is a selection node whose child is n_1 , and its selection variables are $\bar{Y}$, then the adornment of n is obtained from the adornment of n_1 by changing the mapping for the variables in $\bar{Y}$ from f to b . If n is a join node whose children are n_1 (left) and n_2 (right), then the adornment of n satisfies the following constraints: (a) a variable in $adorn(n)$ is mapped to $_$ only if it is mapped to $_$ in both $adorn(n_1)$ and $adorn(n_2)$, (b) if a variable is mapped to b in $adorn(n)$ if it is mapped to b in $adorn(n_1)$ or if it is mapped to b in $adorn(n_2)$ and mapped to $_$ in $adorn(n_1)$, (c) the rest of the variables are bound to f . The set of variables whose values are passed from the left subtree to the right one are those that are mapped to b in $adorn(n_2)$ and mapped to f in $adorn(n_1)$. If this set is non-empty, then the join is a dependent join, otherwise it is just a regular join.

Each of the query execution plans p has an associated cost, denoted by $cost(p)$. Our optimizer includes a component which takes a query execution plan, as described above, and chooses a physical implementation for each of the relational operators. Therefore, the cost of a query execution plan is the cost of the best physical query execution plan implementing it. The cost of atomic query execution plans are directly deduced from the cost associated with the data access descriptions, and the cost of the non-atomic query execution plans is an estimate based on the cost of the operator and the cost of the subplans. The particular cost function we use is orthogonal to the search strategy that our algorithm employs, though it can, in some cases, influence the effectiveness of our pruning methods. We make the monotonicity assumption about our cost model which is a reasonable one in practice: if the plan P' is obtained from the plan P by replacing a subplan P_1 of P by a cheaper and equivalent subplan P_2 , then P' is cheaper than P .

A query execution plan that covers all the conjuncts in the query and whose adornment maps precisely the bound variables in the query to b is called a *complete* query execution plan.

¹Formally, special care needs to be given to the case in which a conjunct $R(\bar{X})$ has the same variable in the positions corresponding to two attributes A and B . In this case, we consider a modification of the binding patterns for R , where whenever either A or B are mapped to b , then both are mapped to b .

Problem definition: Given a set of data access descriptions and a query q , our goal is to find a complete query execution plan for q whose cost is minimal.

Before ending this section we introduce several terms that will be convenient in our discussion. Two query execution plans are considered *equivalent* if they are labeled with the same set of conjuncts from the query and have identical adornments. A query execution plan is *viable* if it can be part of a complete query execution plan. An adornment bp_1 is said to be *weaker* than an adornment bp_2 , denoted by $bp_1 < bp_2$ if every variable that is mapped to b in bp_1 is also mapped to b in bp_2 , and the two adornments map the same set of variables to $_$. A query execution plan p_1 *covers* a query execution p_2 if they are labeled with the same set of conjuncts from the query, and the adornment of p_1 is weaker than the adornment of p_2 .

Example 2.3 Consider the following chain query whose bound variables are x_0 and x_4 :

$$R_1(x_0, x_1), R_2(x_1, x_2), R_3(x_2, x_3), R_4(x_3, x_4), x_0 = c_1, x_4 = c_2$$

where each relation symbol R_i has two adornments: bf and fb . Consider the following equivalence classes:

- The equivalence classes $[R_1(x_0, x_1), R_2(x_1, x_2)]_{\{x_0\}}$ and $[R_1(x_0, x_1), R_2(x_1, x_2)]_{\{x_2\}}$ are valid and viable.
- The equivalence class $[R_1(x_0, x_1), R_3(x_2, x_3), R_4(x_3, x_4)]_{\{x_1, x_2\}}$ is valid but not viable, because a plan that includes only R_2 cannot produce bindings for both required inputs x_1 and x_2 .
- The equivalence class $[R_1(x_0, x_1)]_{\{\}} is not valid.$
- The equivalence class $[R_1(x_0, x_1), R_2(x_1, x_2)]_{\{x_0\}}$ covers the equivalence class $[R_1(x_0, x_1), R_2(x_1, x_2)]_{\{x_0, x_2\}}$.

Intuitively, two equivalent query execution plans solve the same subquery.² Obviously, the equivalence relation partitions the set of query execution plans into classes. All the plans in the same equivalence class are labeled with the same set of conjuncts and the same adornments. In our discussion we refer to coverage and viability of classes of plans, with the following meanings. If a plan is viable, then all the plans in the same equivalence class are viable. If a plan p_1 covers a plan p_2 then any plan p'_1 equivalent to p_1 will cover any plan p'_2 equivalent to p_2 . Intuitively, if P_1 covers P_2 , then by applying selections to any of the plans in P_1 we can obtain plans in P_2 .

3 Optimization with Binding Patterns

In this section we show through a set of examples how some of the basic properties underlying System-R style optimization need to be reconsidered in the presence of binding-pattern limitations. In fact, we will compare not with System-R, but rather with the Garlic data integration system [28] which partially handles binding pattern limitations within the framework of System-R. We begin by arguing that it is essential to consider query execution plans annotated by their input variables. This will be the key feature distinguishing our algorithm from that in Garlic.

“Open” partial query execution plans: in order to compare our approach with that of Garlic, we distinguish two classes of partial query execution plans. An *open* partial query execution plan is a non-atomic plan (i.e., with more than one relation), which *cannot* be executed only with the constants available in the query. Of course, in order for an open subplan to be part of a complete plan, it must receive bindings from some other parts of the plan. A *closed* subplan is one that can be executed given the bindings from the query.

Garlic’s search strategy only considers closed partial query execution plans.³ We now argue that optimal plans may include open subplans, and therefore, it is important to consider a larger space of plans. Subsequently we show that looking at this larger search space has several important ramifications to the properties of our search space.

²In order to simplify the discussion, we ignore the issue of interesting orders produced by query execution plans.

³For atomic plans (plans on a single relation) Garlic considers one plan for every viable binding pattern.

Example 3.1 Consider the following query $[R(x, y), S(y, z), T(z, w)]_{\{y\}}$ and suppose that the only access patterns allowed to the relations R, S and T are: $R(x^f, y^f)$, $S(y^b, z^f)$, $S(y^f, z^b)$, $T(z^f, w^b)$ and $T(z^b, w^f)$, i.e., it is possible to retrieve all the tuples of the relation R , but for the relations S and T we need to give either the value of the first or the second attribute.

If we build the set of partial query execution plans according to the Garlic strategy, we will construct in a bottom-up fashion, for each equivalence class, all the corresponding query execution plans. For example, in the first iteration we will keep the plans for $R(x^f, y^f)$, $S(y^b, z^f)$ (because y could be obtained from R), $S(y^f, z^b)$ (because z could be obtained from T) and $T(z^b, w^f)$ (because z could be obtained from S). In the second iteration, only a single plan will be produced: $(R(x^f, y^f) \bowtie_y S(y^b, z^f))$ since this is the *only* plan of size two which is executable given the constants in the query. Finally, the third iteration will yield a single plan: $(R(x^f, y^f) \bowtie_y S(y^b, z^f)) \bowtie_z T(z^b, w^f)$. If we do consider in the generation phase plans which are temporarily *open*, we would construct in the iteration 2 also the plan $S(y^b, z^f) \bowtie_z T(z^b, w^f)$ which would have lead to another final plan $R(x^f, y^f) \bowtie_y (S(y^b, z^f) \bowtie_z T(z^b, w^f))$.

The search strategy employed by Garlic has two important properties. First, it can be shown that if there *exists* a plan for answering the query, then Garlic will find one, even if it is *not* the optimal plan. Second, if only limited implementations of the dependent join operator are available, then it turns out that for any plan that has an open subplan, there exists a closed subplan that model the *same* execution. For example, if we consider that the dependent joins are implemented using only a nested loop algorithm and that the execution follows the iterator model [12], then the two trees shown previously actually model the same execution, and hence it suffices to consider only the closed plans.

The above two points also highlight the limitations of the Garlic approach, which are of special importance in the context of data integration where some of the sources are on a wide-area network. Nested loop join is far from being the most efficient way of implementing a dependent join. For example, in [13] the authors propose several variants of the known join algorithms in order to implement the dependent join algorithm. In the case of such complicated implementation techniques, the two query execution plans above do not result in the same execution. For example, the first query execution plan above cannot model an execution using caching techniques where every time that the expression $(S(y^b, z^f) \bowtie_z T(z^b, w^f))$ is executed for a given value of y , the output tuples are cached (with the input y). This execution *cannot* be modeled by the first tree, since $(S(y^b, z^f) \bowtie_z T(z^b, w^f))$ does not appear as a subexpression. By considering only closed partial query execution plans Garlic will miss a more efficient plan. In the context of data integration it is essential to consider more sophisticated implementations of dependent joins, and hence we need to develop an algorithm that efficiently explores open subplans.

In what follows we describe several other important properties of the search space in the presence of binding pattern limitations.

Refined equivalence classes: as an immediate consequence of the above discussion is that we need to refine our notion of equivalence classes during our join-order enumeration. System-R style optimizers keep the cheapest query execution plan for every equivalence class, where two plans belong to the same equivalence class if they cover the same set of base relations in the query.⁴ In our new context, it is necessary to annotate every plan also with the set of variables that are required as inputs in addition to the set of base relations that are covered. From now on, the *equivalence class* of a plan is determined by the combination of the subquery solved by a plan and its required input variables.

Valid and viable plans: an immediate consequence of the fact that we must consider refined annotations for our plans is that the number of equivalence classes grows significantly, hence leading to a more expensive search problem. Fortunately, two classes of plans can be pruned early in the search: plans belonging to *invalid* or *non-viable* equivalence classes. An equivalence class is invalid if there is no query execution plan for that class that can be executed given the limitations on access patterns. For example, for the previous query, the equivalence class $[S(y, z), T(z, w)]_{\{w\}}$ is invalid. An equivalence class is not viable if none of the plans in the class can be part of a complete query

⁴To simplify our discussion, we ignore the issue of interesting orders in this paper.

execution plan. In Section 4 we show that pruning these two classes of plans has a dramatic effect on the size of the search space, and in Section 5 we show that validity can be checked as part of the enumeration algorithm, and viability can be checked efficiently.

Two additional properties of our search space are important to understand before we can design an appropriate algorithm: the need to explore bushy trees and the special treatment of selections.

The need for bushy trees: as the following example shows, in the presence of limited access patterns, there are cases where the set of left-linear trees includes *only* plans with cartesian products, while the set of bushy trees *does* contain a query execution plan without cartesian products. Hence, if we want to avoid plans with cartesian products, we must search the space of bushy trees.

Example 3.2 Consider the following query $[P(x, y), R(y, z), S(t, w), T(w, z)]_{\{\}}^{\{\}}$, and suppose that the only access patterns allowed to the relations P, R, S and T are: $P(x^f, y^f)$, $R(y^b, z^f)$, $S(t^f, w^f)$ and $T(w^b, z^f)$. It is easy to note that all the linear query execution plans will include a cartesian product. However, the following bushy-tree does not include a cartesian product: $(P(x^f, y^f) \bowtie_y R(y^b, z^f)) \bowtie_z (S(t^f, w^f) \bowtie_w T(w^b, z^f))$

Recall that in the traditional System-R setting, if the query graph is connected, then the space of left-linear trees necessarily includes a plan without cartesian products. Hence, in that context, the query optimizer can limit its search to left-linear trees without having to use cartesian products.

Placement of selections: in the context of System-R it is possible to decouple the decision on join ordering from the decision on placement of selections. The placement could either be made heuristically by pushing selections as far down as possible in the query tree, in a cost-based fashion in a separate phase [19], or mixing the decision on the join order with the decision on the placement of expensive predicates in a dynamic programming style optimization like in [13]. In our context, since we are considering query execution plans that are annotated by variables that must be given as inputs to the plan, the interaction between the placement of the selection and the join ordering is much more subtle. Specifically, in our search space, the join ordering decision may result in plans that differ only w.r.t. placement of selections. This is illustrated by the following example.

Example 3.3 Consider the query $[R(x, y), S(y, z), T(z, w)]_{\{\}}^{\{\}}$, and suppose that the only access patterns allowed to the relations R, S and T are: $R(x^f, y^f)$, $S(y^b, z^f)$, $S(y^f, z^b)$ and $T(z^f, w^f)$. By combining the pattern $R(x^f, y^f)$ with the pattern $S(y^b, z^f)$ via a dependent join operator, we obtain a plan P_1 for the equivalence class $[R(x, y), S(y, z)]_{\{\}}^{\{\}}$. By combining the pattern $R(x^f, y^f)$ with the pattern $S(y^f, z^b)$ via a join operator, we obtain a plan P_2 for the equivalence class $[R(x, y), S(y, z)]_{\{z\}}^{\{\}}$. It is easy to see that the plan P_1 covers the plan P_2 , i.e., by applying a selection on the variable z in the plan P_1 , we obtain a plan P_3 which is equivalent to P_2 . Hence, a plan without *any* selections turns out to be equivalent to a plan *with* selection, and in our case, it may turn out that P_3 is cheaper than P_2 .

In standard System-R optimization the equivalence class $[R(x, y), S(y, z)]_{\{z\}}^{\{\}}$ would have not been considered *at all* since z is not bound in the original query. In our setting, as shown previously, we have to keep one plan per set of conjuncts and set of bound variables (even if they are not bound in the original query). As a consequence, if we ignore selections, we will be in the situation where we *do* consider a plan for the class $[R(x, y), S(y, z)]_{\{z\}}^{\{\}}$, but the plan we think is optimal for this class (i.e., P_2) is actually not the real optimal one.

To conclude this section, we note that in order to perform optimization in the presence of access pattern limitations, the optimizer must search the space of annotated query plans. The algorithm should avoid invalid plans and prune non-viable plans as early as possible. In order to avoid cartesian products, the algorithm needs to consider bushy trees and not only left-deep trees. Furthermore, special care must be given to placing selections and to detect multiple query plans that result in identical executions. In the next section we analyze the size of the search space sanctioned by the conclusions of this section.

Graph	Qep's	LL Qep's	Pqep's in d.p.	LL Pqep's in d.p.
$R_i(u^f, v^f), i = 1, n$ $\text{bound}(q) = \emptyset$	$\binom{2(n-1)}{n-1} (n-1)!$ [17]	$n!$ [17]	$3^n - 2^{n+1} + 1$ [15, 16]	$n(2^{n-1} - 1)$ [15, 16]
$R_i(u^b, v^f), i = 1, n$ $\text{bound}(q) = \{x_0\}$	$\binom{2(n-1)}{n-1} \frac{1}{n}$	1	$\frac{n^3-n}{6}$ [15, 16]	$n - 1$
$R_i(u^f, v^f)$, for $i \in \{i_1, \dots, i_m\}$, $R_i(u^b, v^f)$ for $i \notin \{i_1, \dots, i_m\}$ $\text{bound}(q) = \emptyset$	$\leq \binom{2(n-1)}{n-1} \frac{(n-1)!}{((\frac{n}{m})!)^m}$	$\frac{n!}{((\frac{n}{m})!)^m}$	$\leq \frac{((\frac{n}{m})^3 + 6(\frac{n}{m})^2 + 5(\frac{n}{m}) + 6)^m}{6^m}$	$\leq m(\frac{n}{m} + 1)^m$
$R_i(u^b, v^f), R_i(u^f, v^b)$, $i = 1, n$ $\text{bound}(q) = \{x_0, x_n\}$	$\binom{2(n-1)}{n-1} \frac{2^n}{n}$	2^n	$O(n^6)$	$O(n^4)$
$R_i(u^b, v^f), R_i(u^f, v^f)$, $i = 1, n$ $\text{bound}(q) = \emptyset$	$\binom{2(n-1)}{n-1} n^n$	n^{n+1}	$O(5.36^n)$	$O(n3.73^n)$

Table 1: Bounds on the size of the search space for a dynamic-programming optimizer in the presence of binding patterns. The query is $q : -R_1(x_0, x_1), R_2(x_1, x_2), \dots, R_n(x_{n-1}, x_n), \mathcal{C}$, and the binding patterns for $R_1, \dots, R_n$ are indicated in each row. The column “Qep” represents all total query execution plans (bushy, with cartesian product); next column restricts the count to left-linear plans (“LL”). The next column represents the partial query execution plans considered by dynamic programming, and the last column the left linear plans.

4 The Size of the Search Space

The size of the space that needs to be searched by a query optimizer employing dynamic programming is relatively well understood [15, 17, 16]. In this section we study the effect on the size of the search space in the presence of access pattern limitations, and the associated need to search the space of annotated query execution plans. We present both an analytical and empirical study. The results of this study also justifies some of the choices we made in designing our query optimization algorithm described in Section 5.

4.1 Theoretical Study

Our study examines the size of two measures: the number of valid complete query execution plans and the number of valid partial query execution plans considered by a dynamic-programming style optimizer. Of course, while the number of query execution plans is usually very large (and that of partial plans even larger), dynamic programming only considers a small fraction of the partial plans. For example, for the case of chain queries with n relations, the number of plans without cartesian products is known to be $\binom{2(n-1)}{n-1} \frac{1}{n}$, while the number of bushy partial plans without cartesian products considered by the dynamic programming algorithm is known to be only $\frac{n^3-n}{6}$ [15]. We note that even though we consider the number of plans explored by a dynamic programming optimizer, the results are of interest even if we were to employ a different paradigm, since dynamic programming is sometimes used as a yardstick for the others. For example, [14] show that a classical rule-based optimizer considers, in general, a strictly higher sized search space than dynamic programming, and present an improved rule-based optimization algorithm whose complexity matches that of dynamic programming.

The results of our analysis are shown in Table 1. The table shows the maximal number of valid query execution plans and valid partial query execution plans generated by the dynamic programming algorithm for the cases in which all bushy trees are considered (columns 1 and 3) and for the case in which only left-linear trees are considered (columns 2 and 4). The formulas include query execution plans that have cartesian products.

We focus our study on chain queries with binary relations:

$$q : -R_1(x_0, x_1), R_2(x_1, x_2), \dots, R_n(x_{n-1}, x_n), \mathcal{C}$$

and consider different combinations of access patterns and different sets of bound variables. Thus, in the first line of the table, all relations have access pattern, $R_i(u^f, v^f)$, for $i = 1, n$, and $\text{bound}(q) = \emptyset$. This line represents the classical case, with no access patterns and no selections, and is for comparison purposes (all entries are taken from the references). Note that since we are counting plans with cartesian products, the numbers in the first row apply to *any* query shape, not just chain queries. In the second line all relations have the access pattern, $R_i(u^b, v^f)$, $i = 1, n$, and $\text{bound}(q) = \{x_0\}$. The third line analyses the transition from line 1 to line 2, by letting the number m of relations with binding pattern ff vary from 1 to n (the other $n - m$ relations have binding pattern bf). We assume here that R_1 is always among these m , i.e. we have the access pattern $R_1(u^f, v^f)$: this guarantees that there always exists a query execution plan, although $\text{bound}(q) = \emptyset$. In line four each of the relations has two binding patterns, $R_i(u^b, v^f)$ and $R_i(u^f, v^b)$, and there are two bound variables: $\text{bound}(q) = \{x_0, x_n\}$ (i.e. we can start either from the left or from the right). Finally, in the last line each relation has both binding patterns bf and ff .

Lines two and three illustrate an example where the complexity of join ordering decreases because of the limited access patterns. Line two represents an extreme case, with a single left-linear solution (namely $(\dots((R_1 \bowtie R_2) \bowtie R_3) \dots \bowtie R_n)$). There are several solutions with bushy trees, basically all ways to parenthesize this expression, but still less than in the classical case (line 1) where, in addition, one could take all permutations. The number of plans considered by the dynamic programming algorithm also decreases dramatically from line one to line two. We remark that the number of bushy trees considered here is the same as that considered in the classical case for plans without cartesian products[15, 16]. The next line refines the analysis by allowing a number m of relations R_i to be ff , the rest being bf . Of course, the exact formulas in each entry depend on which m relation one chooses: the table only shows their maximum values, obtained precisely when the m are chosen equidistantly (i.e. $R_1, R_{\frac{n}{m}+1}, R_{2\frac{n}{m}+1}, \dots$). It is interesting to observe that the formulas in this line coincide with those in line 1 for $m = n$, and with those in line 2 for $m = 1$.

Line four considers an interesting particular case when we can “start at both ends”. All left linear plans are obtained by shuffling a join from the left $R_1 \bowtie R_2 \bowtie R_3 \dots \bowtie R_k$ with one from the right, $R_n \bowtie R_{n-1} \bowtie R_{n-2} \dots \bowtie R_{k+1}$ (e.g. like in $R_1 \bowtie R_2 \bowtie R_n \bowtie R_{n-1} \bowtie R_4 \bowtie R_{n-2} \bowtie \dots$): there are 2^n ways to do that. The complexity here is higher than in line two, but still far less than the case without limited access patterns.

Finally, line five illustrates a case when the complexity increases because of the additional access patterns (both ff and bf). The increase however is still within the same general complexity: it increases from one exponential to a higher exponential, and not to, say, a double exponential. For example (comparing lines 1 and 5) the partial query execution plans considered by the dynamic programming algorithm increased from $O(3^n)$ to $O(5.36^n)$.

The key conclusion we draw from the table is that for some query shapes the presence of limited access patterns significantly reduces the number of valid plans; since the dynamic programming algorithm will discard invalid plans, it will be able to explore a significantly smaller space in these cases. At the same time, in other cases, the complexity actually increases, but the increase stays within the general complexity of join ordering (i.e. exponential).

4.2 Experimental study of the size of the search space

The analytical study provides only upper bounds on the size of the search space, in cases where mathematical analysis is possible. In this section we describe a series of experiments designed to measure the impact of the presence of binding patterns on the size of the search space.

We study the effects of several factors on the size of the space: size of the query, number of variables, shape of the query graph and the type and structure of binding patterns. In our study we consider three measures: (1) the number of complete query execution plans, (2) the number of viable partial query execution plans, and (3) the number of valid but possibly non-viable query execution plans. The real size of the search space is (1), but the complexity of our algorithm is *not* proportional to this number, but to (2). The complexity of our algorithm without the viability test would be proportional to (3). We measure those numbers for both left-linear trees, as well as for bushy trees. Finally, we measure the effect of considering plans with cartesian products.

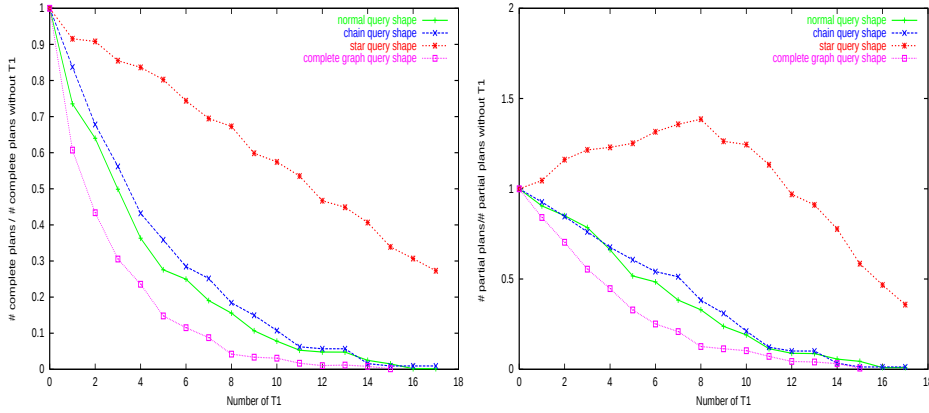


Figure 2: The evolution of the valid partial query execution plans and the complete query execution plans depending on the number of *bind* transformations.

Random query generator: to facilitate the experiments we implemented a random query generator which takes as input: (1) the number of relations, (b) the number of variables and (c) the desired shape of the query graph, and produces as output a conjunctive query with the required properties. We support four kinds of query shapes: chain queries, star queries, complete queries and random. The first three types of queries are well known in the literature [17]. Queries in the last category are constructed as follows: we fix a value of n in advance, the percent of the variables that will appear in 3 relations (i.e., the number of three way joins) and a value m , which will be the percent of variables which belongs to two relations (i.e., the number of two-way joins). The rest of the variables are assumed to belong to a single relation. Given the values n and m , the associations between variables and relation is done randomly. In the rest of the experiments we took the values $n = 1/12$ and $m = 1/3$. Cardinalities of the relations are generated randomly from 1000 to 10000 tuples. The selectivities were randomly chosen between 0.00001 to 1.0. Finally, the cost per tuple associated with each binding pattern was selected randomly between 1 to 1000.

Perturbations on the binding patterns: a simple analysis would show that the size of the search space can be affected by two contradictory factors: (1) when binding pattern limitations become more restrictive, the size of the space decreases, and (2) when new binding patterns are added, the size of the space increases. In order to analyze those two contradictory factors and their respective effect, we apply the following strategy. In both cases, we start from the simple case when all the relations have $fff \dots f$ access. We iteratively apply two transformations on this set of original input binding patterns:

- *bind*: take a binding pattern from the current set, transform one of the f 's into a b , and put it back to the current set.
- *addBinding*: do as in *bind* but do not remove the original binding pattern.

Intuitively, the first transformation would decrease the size of the search space, since the access patterns are more restricted. The second transformation will increase the size of the search space, because there are more ways of accessing the data.

In the figures each point has been obtained from the results of 30 queries generated randomly with the same parameters. We show the average ratios between the number of plans after the transformations and the number of plans for the $ffff$ binding patterns. The results of a more complete set of experiments are given in [18].

The effect of the *bind* transformation: in figure 2 we show how the number of complete query execution plans and the number of partial viable query execution plans vary with the number of applications of the *bind* transformation. The queries have 6 relations, 35 variables, and 12 variables

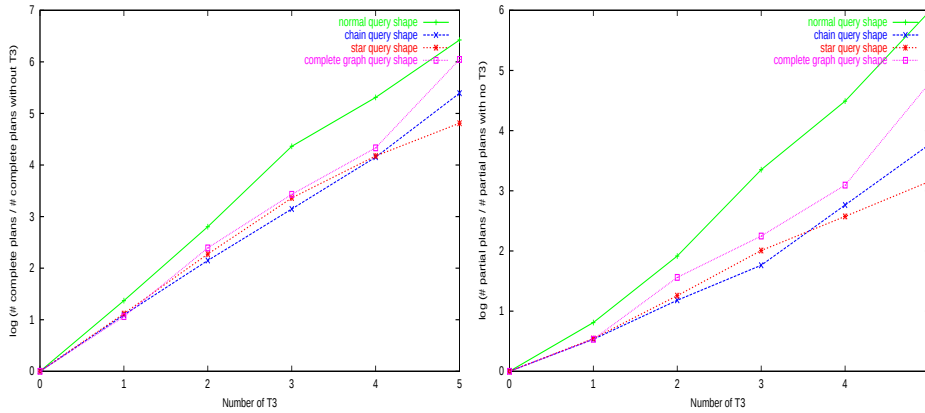


Figure 3: The evolution of the valid partial query execution plans and the complete query execution plans depending on the number of *addBind* transformations.

are bound in the query. We show the results for bushy trees, with cartesian products, and for the four types of queries shapes. We can observe that the size of the search space is decreasing very quickly for all types of queries, as soon as binding patterns are introduced. For example, after 15 applications of the transformation, none of the queries have plans. The number of viable partial query execution plans is globally decreasing, even if sometimes, for star queries, it first slowly increases.

We ran the same tests when varying the other parameters of interest. In particular, we observed that the number of complete query execution plans depends on the shape of the query (it decreases faster for complete queries and much slower for star queries). On the other hand, the size of the query, the shape of the query execution plans (i.e., bushy vs. left-linear) and the consideration of plans with cartesian products does not seem to have a strong effect on the relative average.

The effect of the *addBind* transformation: in Figure 3 we show how the number of complete query execution plan and the number of partial viable query execution plans vary with the number of applications of the *addBind* transformation. Here too the queries have 6 relations, 35 variables and 12 variables are bound. We show the results for bushy trees with cartesian products. The figures show the growth of the search space, on a logarithmic scale.

The four types of queries manifested an exponential growth of the search space depending on the number of the *addBind* transformations. The number of viable partial query execution plans grows accordingly.

The number of non-viable plans: one of the important claims of our paper is that a viability test is essential. In Figure 4 (left) we show how the *total* number of partial plans that can be obtained by a generative algorithm grows when non-viable plans are also considered. As shown by the two bottom curves in the figure, the number of viable partial plans and the number of complete plans are rapidly decreasing. However, the top curve shows that the total number of partial query execution plans increases before it decreases. Hence, this underscores the importance of checking viable plans.

Cartesian products: finally, we tested the effects of allowing plans with cartesian products. In Figure 4 (right) we consider the ratio of the size of the search space (i.e., complete query execution plans) with cartesian product versus the size without cartesian products. We show how this ratio varies with the number of *bind* transformations. The ratio is rather constant or slowly growing when we look at star, complete or normal queries. For chain queries the ratio grows drastically.

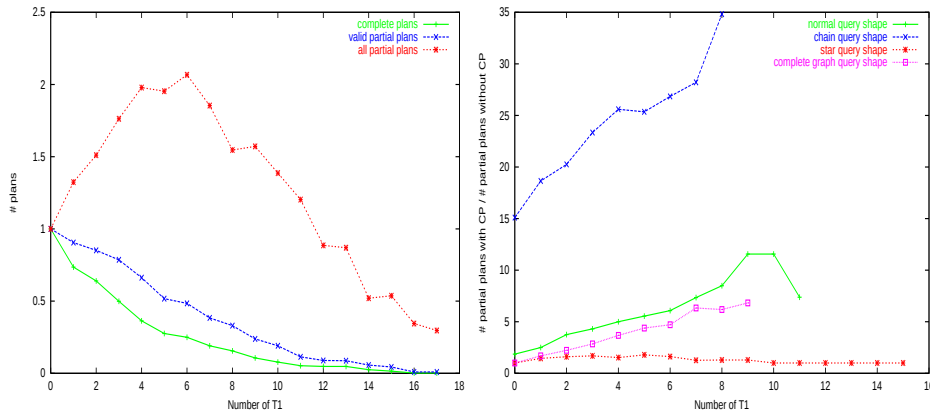


Figure 4: The left graph shows how the number of plans changes with the number of *bind* transformations. Even though the number of viable plans decreases when we perform more *bind* transformations, the number of total valid plans can increase. The right graph shows the effect of including cartesian products in the plans. The growth of the number of plans depends heavily on the shape of the query.

4.3 Consequences of the Size of the Search Space

The analysis of the size of the search space clearly illustrates the need to develop an algorithm that considers only valid and viable plans. The problem that the analysis raises is that it is hard to predict *how* the size of the space will be affected: in some cases, it may be smaller than the traditional case, while in others it can be significantly larger. Hence, in order for an optimization algorithm to be effective in all cases, it must be able to handle large search spaces. The main problem with System-R style optimization when the search space is large is that the *first* plan is produced only towards the end of the optimization. Hence, the approach that we pursue in the next section is to employ a best-first search strategy whose main advantage is to produce a first plan relatively quickly, and improve it as the optimization proceeds.

5 Query optimization algorithm

In this section we describe our query optimization algorithm in detail. We begin by describing the key principles underlying the algorithm, and then focus on some of its important aspects.

5.1 Basic Principles

Our algorithm chooses the optimal plan in the search space characterized by the following properties: (a) bushy trees, (b) plans that include cartesian products, and (c) all the possible placement of selections.

The algorithm is an extension of System-R style optimization, with the following principles:

- At every point in the optimization, the algorithm maintains a set of *partial* query execution plans, S . Each plan $p \in S$ is labeled with the equivalence class to which it belongs and its cost. The equivalence class is specified by the set of conjuncts covered by p and its adornment.
- Initially, the set S contains atomic plans, i.e., plans for accessing a single relation. For a relation R , S contains an atomic query execution plan for every binding pattern describing an access pattern to the tuples of R .
- In the iterative step of the algorithm, we add new plans to S by combining existing plans in S using selection and join operations. We create one resulting plan for every adornment that satisfies the conditions on adornments described in Section 2, hence, creating only valid plans.

- At every point, S contains at most one plan for every equivalence class of query execution plans, which is the cheapest one found thus far.
- The choice of the partial query execution plans to be combined is based on a utility measure. This is significantly different from System-R, where equivalence classes are considered strictly in order of the number of conjuncts they cover.
- In the combination step we prune non-viable plans.

In the following subsections we discuss the main points in which our algorithm differs from System-R: (1) our search strategy, (2) our treatment of the placement of selections, and (3) the detection of useless equivalence classes. The details of the query optimization algorithm are given in Figure 6.

```

Let  $\mathcal{V}$  be the set of variables bound in the query
Let  $\mathcal{C}$  be the empty set
repeat until no change
  if  $e_i(X_i)$  is a conjunct in the query and  $a$  is an adornment for  $e_i$ 
    such that all bound variables in  $a$  are in  $\mathcal{V}$ 
       $\mathcal{V} = \mathcal{V} \cup X_i$ 
       $\mathcal{C} = \mathcal{C} \cup e_i$ 
The query  $Q$  has a plan iff  $\mathcal{C}$  contains all conjuncts in  $Q$ 

```

Figure 5: An algorithm for testing whether a query can be answered given a given set of access patterns to the data. A slight modification of this algorithm yields a test for viability.

5.2 Best First Search

System-R builds query execution plans by considering one equivalence class at a time. The equivalence classes are considered in order of the number of conjuncts they cover. Therefore, the best query execution plan of a class and its cost are determined at one point and are not changed later. The disadvantage of this strategy is that the first complete query execution plan is obtained only at the last phase of the optimization. As the analysis in Section 4 showed, such behavior will not be acceptable in our context.

To address the problem of large search spaces, we employ a best-first search algorithm which *interleaves* the exploration of different equivalence classes. Specifically, we associate a *utility measure* with each partial execution plan we produce. The utility function depends on the number of conjuncts that are covered by the plan and the number of input variables. At each step of the search we choose the partial execution plan with the greatest utility measure, and try to combine it with plans that cover a disjoint set of conjuncts in the query.

The advantage of the best-first search algorithm is that we can tune the utility function to produce a complete plan relatively fast. The main disadvantage of the algorithm is a consequence of the fact that we do not consider each equivalence class in isolation. Therefore, the cost of the best plan for an equivalence class may decrease over time, and hence the cost of the plans using it has to be changed accordingly. The extra bookkeeping resulting from tracking the changes to the costs incurs additional cost. As we show in Section 7 the tradeoff between the two factors is in our favor.

5.3 Placement of selections

Considering all the possible join orderings *and* the possible placement of selections blows up the size of the search space. Previous works have used the powerful heuristic of decoupling the decision about join orders and the placement of selections. In System-R the optimizer heuristically introduces selections as soon as possible in query execution plans (pushing them down as far as possible). In the presence of expensive predicates [19, 20] the placement of the selections is done in a cost-based fashion in a separate phase, *after* the join ordering decision has been made.

As shown in the Example 3.3, it is possible that a plan with a selection as the top operator may be in the same equivalence class as a plan that does not contain *any* selection. Hence, if we

```

Let  $S$  be the set of input binding patterns, each extended to all variables in the query
Test if the query has a query execution plan (validity test). If no, stop.
While new plans can be created:
  Choose  $p_1 \in S$  maximizing the utility measure
  Set  $S'' = \{\}$ 
  Find the set  $S'$  of plans that can be combined with  $p_1$  (i.e., no common conjuncts)
  For each  $p_2 \in S'$  (in the order of their utility measure) do:
    Combine  $p_1$  with  $p_2$  (via a join or a dependent join)  $\rightarrow$  obtain a plan  $p_3$ 
    If there is a  $p_4 \in S$ , such that  $p_4$  covers  $p_3$  and  $cost(\sigma(p_4)) < cost(p_3)$  do  $p_3 = \sigma(p_4)$ 
    Find if there exists a plan  $p_5 \in S$  which is in the same equivalence class as  $p_3$ 
    If yes and the  $cost(p_3) > cost(p_5)$ , ignore  $p_3$ 
    Test if  $p_3$  is viable. If no, ignore  $p_3$ 
    In all the plans  $p_6 \in S$  using  $p_5$ , replace  $p_5$  by  $p_3$  and recalculate the cost of  $p_6$ 
    Add  $p_3$  to  $S$ 
    If there is a  $p_7 \in S$ , such that  $p_3$  covers  $p_7$  and  $cost(\sigma(p_3)) < cost(p_7)$ 
      replace everywhere  $p_7$  by  $\sigma(p_3)$  and recalculate the corresponding cost

If  $S$  contains the equivalence class of the query,
  Then return the optimal plan from  $S$ 
Else
  Let  $S'''$  be the set of plans covering the query
  For each  $p_8 \in S'''$ 
    Generate all the possible placement for the selections in the query not already included in  $p_8$ 
    Choose the optimal among those plans

```

Figure 6: Query optimization algorithm

completely ignore selections during the generation phase, we could miss the optimal plan. The goal of our algorithm is to consider selections in the combination phase only to the extent that it is required in order not to miss optimal plans.

The algorithm considers selections in the following fashion. Suppose we have created a new plan p which is the cheapest one found so far for its equivalence class. Before proceeding, the algorithm checks if it is possible to obtain an equivalent plan to p by applying a selection to a plan that already exists in S . Specifically, the algorithm checks if there exists a plan $p' \in S$, such that p' covers p (i.e., p is equivalent to a selection applied to p'), and the cost of applying the selection to p' is less than the cost of p . In this case, the plan with a selection on p' is added to S instead of p .

Furthermore, the algorithm whether applying a selection on p enables to improve the best plan of another existing equivalent class. Specifically, the algorithm checks whether there exists $p' \in S$ such that a selection on p is equivalent to p' , and the cost of the selection on p is less than the cost of p' . In this case, the algorithm replaces the plan p' in S by the plan with a selection on p .

The effect of the two steps described above is that the set of equivalence classes maintained by S can be characterized as follows: if a class C is in S , then it is there exists at least one query execution plan in C that uses *only* the join operator and no selections. The classes for which all plans contain at least one selection are *not* maintained in S . In a sense, this property entails that the algorithm maintains a *minimal* number of equivalence classes. Furthermore, for these equivalence classes, the algorithm will find the optimal plan with selections.

As a result of the above property, it may be the case that at the end of the generation phase, the equivalence class corresponding to the original query is *not* in S . Assuming the query has at least one query execution plan, this can only happen when all the query execution plans for the query contain selections. In this case, it is easy to check that the optimal plan for the entire query can be obtained by introducing selections in the optimal plans of the equivalence classes covering the query.

Hence, only in the case when the equivalence class of the query does not belong to S at the end of the generation phase, the algorithm applies a second phase, which exhaustively enumerates all the possible placement of selections, but *only* in the optimal plans of the equivalence classes covering the query. Hence, the cost of this phase is relatively small.

5.4 Detection of valid and viable equivalence classes

As stated earlier, our algorithm considers only valid plans and viable query execution plans. The validity of the query execution plans resulting in the combination step is guaranteed by the way we combine plans and generate adornments for the resulting plan.

An algorithm for testing viability of query execution plans can be obtained from by a slight modification of an algorithm for deciding whether a query has any query execution plan. Figure 5 shows a simple inflationary algorithm for deciding whether the a given query has a query execution plan. At every step the algorithm adds to the set of solved subgoals (C) any subgoal whose adornment can somehow be satisfied.

The viability test for a query Q and an equivalence class p for Q is obtained as follows. Let R be a fresh relation symbol. Construct a query Q' by removing from Q all subgoals that appear in p , and replacing them with $R(\bar{X})$, where $\bar{X}$ are p 's variables in p . Associate a single binding pattern with R , namely the adornment of p . It is easy to verify that Q' has a query execution plan if and only if the class p of Q is viable.

6 Implementation

We implemented our algorithm as well as a variant of the dynamic programming algorithm. In order to obtain a fair comparison, we extended the dynamic programming algorithm with a viability test. We use the same data structures (as described shortly) for the two algorithms, and were careful to ensure that the optimizations made in the data structures to efficiently support best-first search do not bias the running times against dynamic programming. The implementation has been done in Java, using JDK 1.0. We briefly explain here some of the choices made when implementing the algorithm we already described.

Data Structure for the Set of Partial Query Execution Plans: A crucial issue that was considered in the implementation is developing a data structure for storing the set of partial plans that have been constructed (denoted by S). An optimal such structure would need to efficiently support the following accesses to the set of plans:

- For a plan p , find all plans q in S , such that p and q have disjoint sets of conjuncts (i.e., the join candidates for p).
- For a plan p , find an equivalent plan p' in S .
- For a plan p , find all the plans q that cover p .
- For a plan p , find all the plans q that are covered by p .

Given these requirements and the observed frequencies of the different accesses, we decided to adapt the following indexing structure for S . Plans are clustered by the set of conjuncts that compose them; note that the join candidates are the same for all the elements of a cluster. In order to avoid repetitive computation of the joinable clusters, the link between joinable clusters is established and materialized when the cluster is given its first member. In addition, the plans in each cluster are indexed by their adornments. It should be emphasized that since equivalent and covering plans belong to the same cluster, and the size of the clusters is relatively small, optimal performance was achieved by *not* adding structures for indexing equivalent and covering plans. Finally, in order to support best-first search, every plan contains a link to the plans using it.

Cost Model: In our experiments we considered a relatively simple cost model. The cost is derived from the cost of the leaf data accesses and standard formulas for computing the cost of joins (hash-join and nested-loop dependent joins). Costs of selections are assumed to be negligible, even though they affect the cardinality of the results. As long as the cost model respects the monotonicity property, the choice of the model is irrelevant to the results we show in the experiments.

Utility Measure: A best-first search algorithm is based on a utility function for choosing the next plan to expand. In our experiments we considered several measures, including (1) the number of conjuncts covered by a plan (2) cost of the plans, (3) number of free variables, and several combinations of the 1–3. Considering only measure (1) resulted in better performance (e.g., up to a factor of 4) in terms of total time and time to first solution, even though the quality of the plans produced early on were not as good as in several of the more complex measures. Considering complex utility measures produces better plans early on in the search but the overhead of the search is significant. Careful tuning of the utility measure is a subject of ongoing research.

7 Experiments

Experiments were ran on a SUN 4 SPARC, under Solaris, using JDK with 100Mb of memory. The main limitation of the algorithm comes from the memory consumption, which is probably increased by the fact that we used Java. Clearly, the use of Java also affects absolute running times of both algorithms, which is why our study focuses on relative numbers. Every point in the graphs is obtained by averaging over 20 queries generated randomly with the same parameters (as described in Section 4). All the experiments are done with queries including 10 relations and 50 variables.

The main motivation for developing this algorithm is that dynamic programming produces the first solution relatively late in the optimization, which yields unacceptable performance in the context of large search spaces. The advantage of best-first search is that it produces the first solutions relatively quickly, but with the added expense in total optimization. Hence, the following two sections quantify the gain in terms of finding the first solution and the price for total optimization time.

7.1 Time to First Solutions

Figure 7 (left) shows the time taken to obtain the first solution for our algorithm and the dynamic programming one. The figure varies the number of *addBind* transformations from 0 to 8. Recall that this variation has the effect of significantly increasing the size of the search space. We observe that the time to first solution for our algorithm is almost constant as the size of the search space increases, while dynamic programming degrades considerably.

Figure 7 (right) shows the ratio between the time to first solution and the total optimization time for both algorithms. Since, as we show below, the total optimization time is in favor of dynamic programming, this graph underscores the superiority of our algorithm w.r.t. finding the first solution. Finally, we observe that the ratio for our algorithm is relatively constant, while it grows for dynamic programming.

It is important to emphasize that our algorithm produces solutions in a relatively steady pace. Hence, we are more likely to obtain a good solution even before dynamic programming produces its first.

7.2 Time for Exhaustive Search

Figure 8 compares the running times for exhaustive search for the two algorithms, as we vary the number of *bind* transformations (left) and as we vary the number of *addBind* transformations (right). We observe that in both cases dynamic programming has a better running time. In the case of *bind* transformations (when the size of the search space decreases) our algorithm takes double time than dynamic programming in the worst case. As the number of *bind* transformation increases, and hence the size of the space decreases, the differences between the running times are negligible. In the case of *addBind* transformations, the running time of both algorithms grows exponentially (note that the Y axis is on a logarithmic scale). Even though our algorithm performs worse, the general growth tendency is the same as for dynamic programming.

In conclusion, we have shown that our algorithm produces first answers considerably faster than dynamic programming. In cases when the search space is relatively small, the additional price paid by our algorithm is not significant. Finally, we argue that in the cases where we do much worse than dynamic programming are anyway cases in which dynamic programming is not a viable strategy.

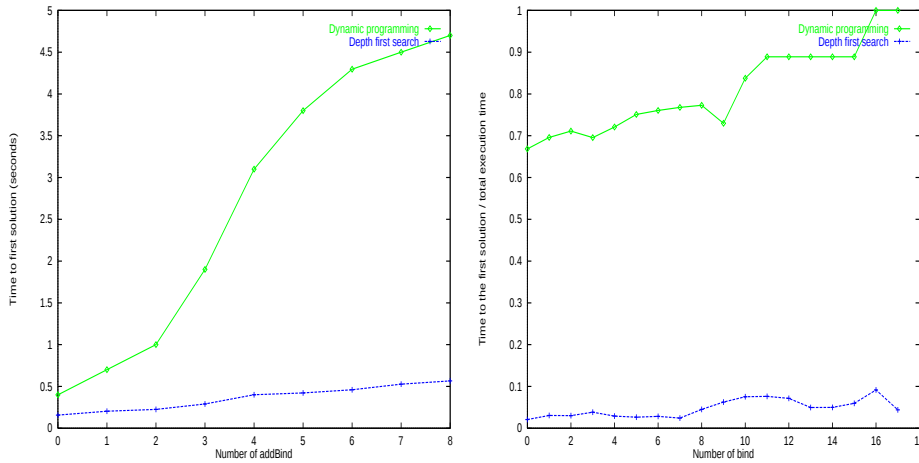


Figure 7: The left graph shows the absolute time taken to find the first solution. The graph on the right shows the ratio between the time to first solution and the time for exhaustive search.

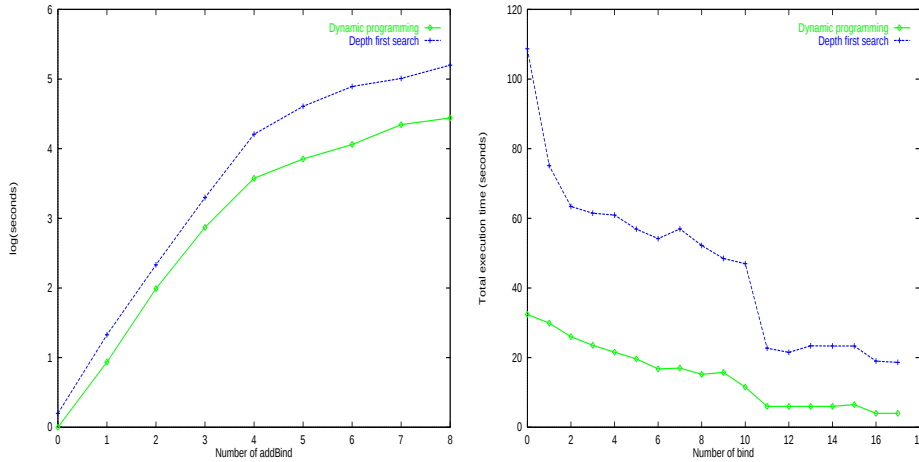


Figure 8: The graphs show the time for exhaustive search in the case of varying the number of *bind* transformations (left) and varying the number of *addBind* transformations (right).

8 Discussion and Related Work

We described a query optimization algorithm which extends System-R style optimization to accommodate limited access patterns to the data. Our algorithm has several important features that are necessitated by the results of our analysis of the properties of the search space arising in the presence of limited access patterns, and a theoretical and experimental study of the size of the space. In particular, our algorithm searches the space of annotated query plans, and prunes as early as possible in the search plans that are invalid or are not viable. Furthermore, to perform well when the search space is large, the algorithm employs a best-first search strategy to produce a complete plan early in the optimization process. The algorithm also handles the placement of selections in a way that is tailored to this new context. Finally, our experimental results show that the algorithm presents a promising tradeoff between obtaining complete plans early on in the search and time for exhaustive search of the space.

In this paper we focused on extending System-R dynamic programming style optimization to handle limited access patterns. A natural question is whether one of the other query optimization paradigms such as the transformational or randomized approach would be more appropriate. For example, in a transformation-based approach [21] the optimizer would start with *some* initial complete plan, and apply transformations to it in order to find an optimal plan. However, this approach requires a set of transformation rules that take one valid plan into another. In our context, the clas-

sical transformation rules such as associativity and commutativity of joins do not have this property, and therefore applying this approach is quite a bit more subtle. The situation is even worse for a randomized approach, because we cannot be guaranteed to cover all the valid plans. Finally, the analysis of the search problem that we provided in this paper can facilitate future attempts to apply different search paradigms.

As we stated in the introduction, one of the reasons for the existence of limited access patterns is the mismatch between the logical and physical views of the data. In our work, binding patterns were used to describe such mismatches. Tsatalos et al. [22] describe GMAPs which are also a mechanism for describing different storage patterns of the data. Using GMAPs, one can describe storage structures in which the stored data is a result of projections, selections and joins on the logical schema of the data. For example, using GMAPs it is possible to describe secondary indexes, path indexes and field replication. GMAPs and binding patterns characterize disjoint sets of mismatches between the logical and physical views of the data. To combine the two families of mismatches, we need to extend the algorithm of in [22] in several ways. First, as we did here, we need to consider annotated query execution plans. Note that in [22] the execution plans manipulate GMAPs (which can be thought of as materialized views) rather than database relations. Second, the join enumeration algorithm needs to consider a plans of larger size. It follows from [26] that in the combined context of binding patterns and GMAPs the query execution plan may require *more* joins than the number of relations in the query, and hence a relation mentioned once in the query may appear in more than one leaf in the query execution plan. In particular, given a query with n subgoals and m variables, the optimization algorithm needs to consider all the plans with $n + m - 1$ joins in order to be guaranteed to find a plan if one exists. An extension of GMAP to deal with binding pattern limitations would also be useful for producing plans that include the novel scan operator proposed by Cluet and Moerkotte in the context of semi-structured data [23].

Stocker and Kossman [24] propose *iterative dynamic programming*, as an alternative method for dealing with unpredictably larger search spaces (which may be due to causes other than binding patterns). In contrast to our method, they begin with pure dynamic programming, and adapt it later if the search space proves to be too large. The challenge in their approach is to detect when to switch from pure dynamic programming to their more directed method.

As we noted the problem of limited access to stored data also arises in the context of data integration. For example, if we are accessing data from a web site, we may only be able to query it with certain binding patterns. Hence, the problem of building query execution plans when only limited access patterns are available has been considered in work on data integration [25, 26, 27]. However, in that work they addressed the question of whether there *exists* some ordering of accesses to the data sources such that an answer to the query can be obtained. The question of finding an *optimal* order was not considered. A related issue is query optimization when the *capabilities* of the data sources are varying. Haas et al. [28] consider query optimization in the context of the Garlic system, where each data source may have different capabilities for performing joins internally. Vassalos and Papakonstantinou describe a powerful language for describing source capabilities [29].

This work is part of a bigger effort to build a query optimizer for contexts in which there is a mismatch between the logical and physical views of the data. As we already mentioned, one future direction is building an optimizer that can support both GMAPs and binding patterns. A second direction is to extend our optimizer to explore query execution plans that are directed acyclic graphs rather than trees. It has already been noted that such plans are useful even in the traditional optimization context, but this is even more so in the presence of limitations on access patterns.

References

- [1] Berthold Reinwald and Hamid Pirahesh. SQL Open Heterogeneous Data Access. In *Proceedings of the ACM SIGMOD International Conference on Management of Data*, 1998.
- [2] Serge Abiteboul. Querying Semi-Structured Data. In *International Conference on Database Theory*, Delphi, Greece, 1997.
- [3] Peter Buneman. Semistructured Data. In *Proc. of the ACM SIGACT-SIGMOD-SIGART Symposium on Principles of Database Systems (PODS)*, Tucson, Arizona, pages 117–121, 1997.
- [4] Alin Deutsch, Mary Fernandez, Daniela Florescu, Alon Levy and Dan Suciu. A Query Language for XML. <http://www.research.att.com/~mff/xml/w3c-note.html>, 1998.

- [5] Odysseas G. Tsatalos, Marvin H. Solomon and Yannis E. Ioannidis. The GMAP: A Versatile Tool for Physical Data Independence. In *Proceedings of the 20th International Conference on Very Large Databases*, Santiago, Chile, pages 367–378, 1997.
- [6] Jeffrey D. Ullman. Information Integration using Logical Views. In *International Conference on Database Theory*, Delphi, Greece, 1997.
- [7] H. Garcia-Molina, Y. Papakonstantinou, D. Quass, A. Rajaraman, Y. Sagiv, J. Ullman and J. Widom. The TSIMMIS Project: Integration of Heterogeneous Information Sources. In *Journal of Intelligent Information Systems*, vol. 8(2):117-132, March 1997.
- [8] Praveen Seshadri, Joseph M. Hellerstein, Hamid Pirahesh, T. Y. Cliff Leung, Raghu Ramakrishnan, Divesh Srivastava, Peter J. Stuckey and S. Sudarshan. Cost-Based Optimization for Magic: Algebra and Implementation. In *Proceedings of the ACM SIGMOD International Conference on Management of Data*, pages 435-446, 1996.
- [9] Raghu Ramakrishnan and Jeffrey D. Ullman. A survey of deductive database systems. In *Journal of Logic Programming*, 23(2):125-149, 1995.
- [10] Surajit Chaudhuri and Kyuoseok Shim. Query Optimization in the presence of Foreign Functions. In *Proceedings of the 19th International Conference on Very Large Databases*, 1993.
- [11] Surajit Chaudhuri, Umeshwar Dayal and Tak Yan. Join Queries with External Text Sources: Execution and Optimization Techniques. In *Proceedings of the ACM SIGMOD International Conference on Management of Data*, 1995.
- [12] Gotz Graefe. Query Evaluation Techniques for Large Databases. In *ACM Computing Surveys*, 25(2):73–170, June 1993.
- [13] Surajit Chaudhuri and Kyuoseok Shim. Optimization in the presence of user-defined predicates. In *Proceedings of the 22nd International Conference on Very Large Databases*, 1996.
- [14] A.Pellenkoft, C.Galindo-Legaria and M.Kersten. The complexity of tranformation-based join enumeration. In *Proceedings of the 24th International Conference on Very Large Databases*, pages 306-315, 1998.
- [15] K.Ono and G.Lohman. Measuring the complexity of join enumeration in query optimization. In *Proceedings of the 16th International Conference on Very Large Databases*, 1990.
- [16] W.Scheufele and G.Moerkotte. Efficient dynamic programming algorithms for ordering expensive joins and selections. In *Proceedings of the 8th International Conference on Extending Database Technology*, 1998.
- [17] M.Steinbrunn, G.Moerkotte and A.Kemper. Heuristic and Randomized Optimization for the Join. In *Proceedings of the 23rd International Conference on Very Large Databases Journal*, 6(3):191-208, 1997.
- [18] Daniela Florescu, Alon Levy, Ioana Manolescu and Dan Suciu. Query Optimization in the Presence of Limited Access Patterns, Technical Report, INRIA, 1999.
- [19] Joseph M. Hellerstein and Michael Stonebraker. Predicate Migration: Optimizing Queries with Expensive Predicates. In *Proceedings of the ACM SIGMOD International Conference on Management of Data*, pages 267-276, 1993.
- [20] Joseph M. Hellerstein. Practical Predicate Placement. In *Proceedings of the ACM SIGMOD International Conference on Management of Data*, pages 325-335, 1994.
- [21] Goetz Graefe. Volcano - An Extensible and Parallel Query Evaluation System. In *IEEE Transactions on Data and Knowledge Engineering*, 6(1):120-135, 1994.
- [22] Odysseas G. Tsatalos, Marvin H. Solomon and Yannis E. Ioannidis. The GMAP: A Versatile Tool for Physical Data Independence. In *International Conference on Very Large Databases Journal*, 5(2):101–118, 1996.
- [23] Sophie Cluet and Guido Moerkotte. Query Processing in schemaless and semistructured context, unpublished manuscript, 1998.
- [24] Konrad Stoker and Donald Kossmann. Iterative Dynamic Programming: A New Class of Query Optimization Algorithms. Submitted for publication, 1998
- [25] K. A. Morris. An algorithm for ordering subgoals in NAIL!. In *Proc. of the ACM SIGACT-SIGMOD-SIGART Symposium on Principles of Database Systems(PODS)*, Chicago, Illinois, pages 82–88, 1988.
- [26] Anand Rajaraman, Yehoshua Sagiv and Jeffrey D. Ullman. Answering Queries Using Templates with Binding Patterns. In *Proc. of the ACM SIGACT-SIGMOD-SIGART Symposium on Principles of Database Systems(PODS)*, San Jose, CA, 1995.
- [27] Alon Y. Levy, Anand Rajaraman and Joann J. Ordille, Querying Heterogeneous Information Sources Using Source Descriptions. In *Proceedings of the 22nd International Conference on Very Large Databases*, Bombay, India, 1996.

- [28] Laura Haas, Donald Kossmann, Edward Wimmers and Jun Yang. Optimizing Queries Across Diverse Data Sources. In *Proceedings of the 23rd International Conference on Very Large Databases*, Athens, Greece, 1997.
- [29] Vasilis Vassalos and Yannis Papakonstantinou. Describing and Using the Query Capabilities of Heterogeneous Sources. In *Proceedings of the 23rd International Conference on Very Large Databases*, Athens, Greece, 1997.

9 Appendix

The first row contains known results for optimization of join-ordering without binding patterns, for comparison purposes. Note that since we allow qep's with Cartesian products all formulas in this row apply to *any* query graph, not just a chain.

The second row considers a standard chain query with binding patterns. There exists a single valid permutation, hence the number of qep's is the number of ways to parenthesize $R_1, \dots, R_n$, and there exists a single LL qep. The number of partial qep's considered by the dynamic programming algorithm is the same as the number of plans without binding patterns and without Cartesian product, for which [15, 16] provide the analysis.

Valid permutations Most of our analysis rely on the following observation. Let $R_1, \dots, R_n$ be all goals involved in some query, and let $R_{\sigma(1)}, \dots, R_{\sigma(n)}$ be some permutation of these goals. Then there exists some qep p whose leaves enumerate the goals in this particular order iff there exists a left-linear qep with the same permutation. The proof of this property relies on simple associativity properties for joins and dependent joins. We say that the permutation $R_{\sigma(1)}, \dots, R_{\sigma(n)}$ is **valid** if there exists some qep whose leaves enumerate the goals in that order. Checking if a permutation is valid is easy: simply run the inflationary validity algorithm with the restrictions that goals are considered precisely in the order of the given permutation. We use this property in an essential way during our analysis. For instance in order to compute the number of valid qep's it suffices to compute the number of valid permutations, say $P(n)$, then to multiply by $\binom{2(n-1)}{n-1} \frac{1}{n}$ (the number of ways to parenthesize a n -way join): of course, the number of LL qep's is exactly $P(n)$. Note that this property holds only because of our assumption that qep's are allowed to have Cartesian products, and fails otherwise.

The third row considers a chain query in which we can start at both ends. Here any valid permutation is obtained by merging $R_1, R_2, \dots, R_k$ with $R_n, R_{n-1}, \dots, R_{k+1}$, for any $k = 0, n$. Hence there are 2^n valid permutations. The dynamic programming algorithm considers qep's for equivalence classes of the form $\{R_i^{bf}, R_{i+1}^{bf}, \dots, R_{j-1}^{bf}, R_k^{fb}, R_{k+1}^{fb}, \dots, R_l^{fb}\}$, where $1 \leq i \leq j < k \leq l \leq n$. I.e. each equivalence class has a “left-to-right” chain and a “right-to-left” chain. For each such equivalence class the d.p. algorithm considers all ways of splitting the left-to-right chain and the right-to-left chain into two subchains: each choice can be characterized by choosing two indexes x, y s.t. $i - 1 \leq x \leq j$ and $k - 1 \leq y \leq l$, except for the two choices $x = i - 1 \wedge y = k - 1$ and $x = j \wedge y = l$ (which yield a left empty plan, or a right empty plan respectively). There are $\approx \binom{n}{6}$ ways to choose the six indexes, hence $O(n^6)$. For LL qep's, it suffices to restrict x and y to $j - 1$ and $l - 1$ respectively, hence $\approx \binom{n}{4} = O(n^4)$.

The fourth line considers the case when m of the relations $R_1, \dots, R_n$ have binding patterns ff instead of bf : to have at least one valid plan, we must have R_1 among those m relations. Here our analysis starts by observing that the remaining $m - 1$ relations divide the chain into m subchains of length $k_1, \dots, k_m$, s.t. $\sum_i k_i = n$. Then each valid permutation is obtained by shuffling (merging) the m chains of lengths $k_1, k_2, \dots, k_m$ respectively: hence there are $\frac{n!}{k_1! \dots k_m!}$ valid permutations. This number maximizes when $k_1 = k_2 = \dots = k_m = \frac{n}{m}$, hence the number of qep's and of LL qep's. For the number of partial qep's considered by d.p., we observe that each plan constructed by the d.p. algorithm consists of m pairs of adjacent subchains. Each of the subchains may be empty, the only prohibited cases are when *all* left subchains are empty or when *all* right subchains are empty. Considering only the first chain (of length k_1), for each $i = 1, k_1$ there are $k_1 - i + 1$ chains of length i , and each can be split in $i + 1$ ways into two adjacent chains (including left empty chain and right empty chain). Hence there are $1 + \sum_{i=1}^{k_1} (k_1 - i + 1)(i + 1)$ total pairs of adjacent subchains (the 1 stands for both subchains being empty). This sum is $\frac{k_1^3 + 6k_1^2 + 5k_1 + 6}{6}$. Considering now all m chains, the number of plans constructed by the d.p. algorithm is bounded by $f(k_1)f(k_2) \dots f(k_m)$,

where $f(x) = \frac{x^3+6x^2+5x+6}{6}$. The product also includes the prohibited cases (hence it is only an upper bound). It maximizes when $k_1 = k_2 = \dots = k_m = \frac{n}{m}$, because the function $g(x) \stackrel{\text{def}}{=} \ln(f(x))$ is concave⁵ hence the table entry. Note that this is asymptotically equal to the entry in the first row ($O(3^n)$), when $m = n$, and to that in the second row ($O(n^3)$), when $m = 1$. The classical case (row 1), with n perturbations, represents therefore a worst case for the d.p. algorithm.

For left-linear tree, the analysis is simpler. A partial plan consists of *initial* segments of the m chains, and each such plan can be extended in at most m ways. Hence the number of LL pqep's is bounded by:

$$m \prod_{j=1}^m (k_j + 1) \leq m \left(\frac{n}{m} + 1\right)^m$$

Again, the classical case turns out to be the worst case.

The fifth row represents a worst case for optimization with binding patterns: for each goal we can choose either ff or bf . In order to count the number of valid permutations, consider the case when exactly m relations are ff (R_1 must be among them) which divide the chain into m nonempty subchains of length $k_1, \dots, k_m$ (i.e. the m relations are $R_1, R_{1+k_1}, R_{1+k_1+k_2}, \dots$ and $\sum_{i=1}^m k_i = n$). All valid permutations with these binding patterns are obtained from merging the m chains $(R_1, R_2, \dots, R_{k_1}), (R_{1+k_1}, \dots, R_{1+k_1+k_2}), \dots$, hence there are $\frac{n!}{k_1! \dots k_m!}$ such valid permutations. Summing up over all choices of m and $k_1, \dots, k_m$, we get the following upper bound on the total number of valid permutations:

$$\leq 1^n + 2^n + 3^n + \dots + n^n \leq n^{n+1}$$

(This bound overshoots the real number because it also includes choices for $k_i = 0$ corresponding to empty chains). For the number of partial plans explored by the d.p. algorithm we proceed as in row four: for each choice of $1 \leq m \leq n$ and $k_1, \dots, k_m \geq 1$ s.t. $k_1 + \dots + k_m = n$ the dynamic programming algorithm generates $f(k_1)f(k_2)\dots f(k_m)$ partial plans. Summing up we get the following upper bound for the total number of partial plans:

$$\sum_{m \geq 0, k_1, \dots, k_m \geq 0, k_1 + \dots + k_m = n} f(k_1)f(k_2)\dots f(k_m) \stackrel{\text{def}}{=} S(n)$$

For $S(n)$ we derive the following recurrence:

$$\begin{aligned} S(0) &= 1 \\ S(n) &= f(n)S(0) + f(n-1)S(1) + \dots + f(1)S(n-1) \end{aligned}$$

Define $G(z) = \sum_{n \geq 0} S(n)z^n$ and $H(z) = \sum_{n \geq 0} f(n+1)z^n$ the generating functions for $S(n)$ and $f(n)$, we have:

$$\begin{aligned} G(z) &= 1 + zG(z)H(z) \\ G(z) &= \frac{1}{1 - zH(z)} \end{aligned}$$

Hence we have to (1) compute $H(z)$, (2) then compute $G(z)$. Before doing an exact analysis we do a qualitative discussion. First $H(z)$ is of the form $\frac{\varphi(z)}{(1-z)^4}$ where $\varphi(z)$ is a polynomial of degree 3 (this is because $f(n)$ is a polynomial of degree 3). Hence $G(z) = \frac{(1-z)^4}{\psi(z)}$ with $\psi(z) = (1-z)^4 - z\varphi(z)$ a polynomial of degree 4. Assuming $\psi(z)$ has four isolated roots, C_1, C_2, C_3, C_4 , (i.e. no multiple roots) we have:

$$G(z) = A_0 + \frac{A_1}{z - C_1} + \frac{A_2}{z - C_2} + \frac{A_3}{z - C_3} + \frac{A_4}{z - C_4}$$

which gives us $S(n) = A_1 \frac{1}{C_1^{n-1}} + A_2 \frac{1}{C_2^{n-1}} + A_3 \frac{1}{C_3^{n-1}} + A_4 \frac{1}{C_4^{n-1}}$ for $n > 0$ ($S(0)$ has the additional term A_0), hence $S(n) = O(C^n)$ where C is the largest reciprocal of all four roots. To actually determine

⁵A function $g(x)$ is concave on some interval $[a, b]$ when $g''(x) < 0, \forall x \in [a, b]$. Our function g is concave: indeed we have $g''(x) = (\ln(f(x)))'' = \frac{f''(x)f(x) - (f'(x))^2}{f^2(x)} = \frac{-6(5x^4 + 12x^3 + 144x^2 + 120x - 11)}{(f(x))^2} \leq 0$ for $x \in [1, \infty)$. A concave function has the property $\frac{g(k_1) + g(k_2) + \dots + g(k_m)}{m} \leq g(\frac{k_1 + \dots + k_m}{m})$.

C we perform next an exact quantitative analysis. To compute $H(z)$ easier we notice that $f(n)$ can be written as:

$$f(n) = \frac{n(n+1)(n+2)}{6} + \frac{n(n+1)}{2} + 1$$

hence we get:

$$\begin{aligned} H(z) &= \sum_{n \geq 0} f(n+1)z^n \\ &= \sum_{n \geq 0} \left(\frac{(n+1)(n+2)(n+3)}{6} z^n + \frac{(n+1)(n+2)}{2} z^n + z^n \right) \\ &= \frac{1}{6} (\sum z^{n+3})''' + \frac{1}{2} (\sum z^{n+2})'' + \sum z^n \\ &= \frac{1}{6} \left(\frac{1}{1-z} \right)''' + \frac{1}{2} \left(\frac{1}{1-z} \right)'' + \frac{1}{1-z} \\ &= \frac{1}{(1-z)^4} + \frac{1}{(1-z)^3} + \frac{1}{1-z} \\ &= \frac{-z^3 + 3z^2 - 4z + 3}{(1-z)^4} \\ G(z) &= \frac{1}{1 - zH(z)} \\ &= \frac{(1-z)^4}{2z^4 - 7z^3 + 10z^2 - 7z + 1} \end{aligned}$$

The denominator has the following roots: $C_1 = 0.8376 + 0.9676i$, $C_2 = 0.8376 - 0.9676i$, $C_3 = 1.6384$, $C_4 = 0.1863$. The largest reciprocal (and the only one greater than 1) is $C = 1/C_4 = 5.3668$.

We briefly present next how to derive the number of LL plans considered by the dynamic programming algorithm, $T_{LL}(n)$. Denoting $f_{LL}(x) = (x+1)$, and reasoning as above we deduce:

$$T_{LL}(n) = \sum_{m \geq 1, k_1, \dots, k_m \geq 1, k_1 + \dots + k_m = n} m \Pi_{i=1}^m f_{LL}(k_i)$$

Let us also denote:

$$S_{LL}(n) = \sum_{m \geq 1, k_1, \dots, k_m \geq 1, k_1 + \dots + k_m = n} \Pi_{i=1}^m f_{LL}(k_i)$$

(i.e. without the factor m). We derive the following recurrences:

$$\begin{aligned} S_{LL}(0) &= 1 \\ S_{LL}(n) &= f_{LL}(n)S_{LL}(0) + f_{LL}(n-1)S_{LL}(1) + \dots + f_{LL}(1)S_{LL}(n-1) \\ T_{LL}(0) &= 0 \\ T_{LL}(n) &= f_{LL}(n)(T_{LL}(0) + S_{LL}(0)) + f_{LL}(n-1)(T_{LL}(1) + S_{LL}(1)) + \dots + f_{LL}(1)(T_{LL}(n-1) + S_{LL}(n-1)) \end{aligned}$$

Next denote $H_{LL}(z)$, $G_{LL}(z)$, $K_{LL}(z)$ the generating functions:

$$\begin{aligned} H_{LL}(z) &\stackrel{\text{def}}{=} \sum_{n \geq 0} f_{LL}(n+1)z^n \\ G_{LL}(z) &\stackrel{\text{def}}{=} \sum_{n \geq 0} S_{LL}(n)z^n \\ K_{LL}(z) &\stackrel{\text{def}}{=} \sum_{n \geq 0} T_{LL}(n)z^n \end{aligned}$$

and derive the following relations from the recurrences:

$$G_{LL}(z) = 1 + zH_{LL}(z)G_{LL}(z)$$

$$\begin{aligned}
K_{LL}(z) &= zH_{LL}(z)(G_{LL}(z) + K_{LL}(z)) \\
G_{LL}(z) &= \frac{1}{1 - zH_{LL}(z)} \\
K_{LL}(z) &= \frac{zH_{LL}(z)G_{LL}(z)}{1 - zH_{LL}(z)} \\
&= \frac{zH(z)}{(1 - zH(z))^2}
\end{aligned}$$

Compute $H_{LL}(z)$:

$$\begin{aligned}
H_{LL}(z) &= \sum_{n \geq 0} (n+2)z^n = \frac{1}{z} \sum_{n \geq 0} (n+2)z^{n+1} \\
&= \frac{1}{z} \left(\sum_{n \geq 0} z^{n+2} \right)' = \frac{1}{z} \left(\frac{1}{1-z} - 1 - z \right)' \\
&= \frac{1}{z} \left(\frac{1}{(1-z)^2} - 1 \right)
\end{aligned}$$

Sustituting in $K_{LL}(z)$ gives us:

$$K_{LL}(z) = \frac{z(2-z)(1-z)^2}{(z^2 - 4z + 1)^2}$$

The denominator has two double roots, $C_1 = 3.7321$ and $C_2 = 0.2679$, henc $T_{LL}(n) = A_1 \frac{1}{C_1^n} + A_2 \frac{n}{C_1^{n-1}} + A_3 \frac{1}{C_2^n} + A_4 \frac{n}{C_2^{n-1}}$ for $n \geq 1$ ($T_{LL}(0)$ is also influenced by H_{LL} 's numerator). Since C_1 has the largest reciprocal, 3.7321, we get $T_{LL}(n) = O(n3.73^n)$.



Unité de recherche INRIA Lorraine, Technopôle de Nancy-Brabois, Campus scientifique,
615 rue du Jardin Botanique, BP 101, 54600 VILLERS LÈS NANCY
Unité de recherche INRIA Rennes, Irisa, Campus universitaire de Beaulieu, 35042 RENNES Cedex
Unité de recherche INRIA Rhône-Alpes, 655, avenue de l'Europe, 38330 MONTBONNOT ST MARTIN
Unité de recherche INRIA Rocquencourt, Domaine de Voluceau, Rocquencourt, BP 105, 78153 LE CHESNAY Cedex
Unité de recherche INRIA Sophia-Antipolis, 2004 route des Lucioles, BP 93, 06902 SOPHIA-ANTIPOLIS Cedex

Éditeur
INRIA, Domaine de Voluceau, Rocquencourt, BP 105, 78153 LE CHESNAY Cedex (France)
<http://www.inria.fr>
ISSN 0249-6399