

Efficient Query Evaluation on Probabilistic Databases

Nilesh Dalvi and Dan Suciu

April 4, 2004

Abstract

We describe a system that supports arbitrarily complex SQL queries with “uncertain” predicates. The query semantics is based on a probabilistic model and the results are ranked, much like in Information Retrieval. Our main focus is query evaluation. We describe an optimization algorithm that can compute efficiently most queries. We show, however, that the data complexity of some queries is $\#P$ -complete, which implies that these queries do not admit any efficient evaluation methods. For these queries we describe both an approximation algorithm and a Monte-Carlo simulation algorithm.

1 Introduction

Databases and Information Retrieval [5] have taken two philosophically different approaches to queries. In databases SQL queries have a rich structure and a precise semantics. This makes it possible for users to formulate complex queries and for systems to apply complex optimizations, but users need to have a pretty detailed knowledge of the database in order to formulate queries. For example, a single misspelling of a constant in the WHERE clause leads to an empty set of answers, frustrating casual users. By contrast, a query in Information Retrieval (IR) is just a set of keywords and is easy for casual users to formulate. IR queries offer two important features that are missing in databases: the results are *ranked* and the matches may be *uncertain*, i.e. the answer may include documents that do not match all the keywords in the query¹. While several proposals exist for extending SQL with uncertain matches and ranked results [3, 17, 14], they are either restricted to a single table, or, when they handle join queries, adopt an ad-hoc semantics.

To illustrate the point consider the following structurally rich query, asking for an actor whose name is like ‘Kevin’ and whose first ‘successful’ movie appeared in 1995:

¹Some IR systems only return documents that contain all keywords, but this is a feature specific to those systems, and not of the underlying vector model used in IR.

```

SELECT      *
FROM        Actor A
WHERE       A.name  $\approx$  'Kevin'
    and     1995 =
    SELECT  MIN(F.year)
FROM        Film F, Casts C
WHERE       C.filmid = F.filmid
    and     C.actorid = A.actorid
    and     F.rating  $\approx$  "high"

```

Figure 1: An Approximate Query

The two $\approx$ operators indicate which predicates we intend as uncertain matches. Techniques like edit distances, ontology-based distances [13], IDF-similarity and QF-similarity [3] can be applied to a *single* table: to rank all **Actor** tuples (according to how well they match the first uncertain predicate), and to rank all **Film** tuples. But it is unclear how to rank the *entire* query. To date, no system combines structurally rich SQL queries with uncertain predicates and ranked results.

In this paper we propose such a system. We introduce a new semantics for database queries that supports uncertain matches and ranked results, by combining the probabilistic relational algebra [11] and models for belief [4]. Given a SQL query with uncertain predicates, we start by assigning a probability to each tuple in the input database according to how well it matches the uncertain predicates. Then we derive a probability for each tuple in the answer, and this determines the output ranking.

An important characteristic of our approach is that *any* SQL query has a meaning, including queries with joins, nested sub-queries, aggregates, group-by, and existential/universal quantifiers². Queries have now a probabilistic semantics, which is simple and easy to understand by both users and implementors.

The main problem is query evaluation, and this is the focus of our paper. Our approach is to represent SQL queries in an algebra, and modify the operators to compute the probabilities of each output tuple. This is called *extensional semantics* in [11], and is quite efficient. While this sounds simple, the problem is that it doesn't work: the probabilities computed this way are wrong in most cases, and lead to incorrect ranking. In [11], the workaround is to use an *intensional semantics*, which is much more complex and, as we show here, impractical. Our approach is different: we rewrite the query plans, searching for one where the extensional evaluation is correct. We show however that certain queries have a #P-complete data complexity under probabilistic semantics, and hence do not admit a correct extensional plan. While they are not frequent in practice (only 2 out of the 10 TPC/H queries fall in this category, and only when all their predicates are uncertain), we describe two techniques to address

²In this paper we restrict our discussion to SQL queries whose normal semantics is a set, not a bag or an ordered list.

	S^p					
	A	B				
s_1	'm'	1	0.8			
s_2	'n'	1	0.5			

	T^p					
	C	D				
t_1	1	'p'	0.6			

Figure 2: A probabilistic database D^p

$pwd(D^p) =$			
world	prob.	$q(u) : S^p(x, y), T^p(z, u), y = z$	
$D_1 = \{s_1, s_2, t_1\}$	0.24	(b)	
$D_2 = \{s_1, t_1\}$	0.24		
$D_3 = \{s_2, t_1\}$	0.06		
$D_4 = \{t_1\}$	0.06		
$D_5 = \{s_1, s_2\}$	0.16	$q^{pwd}(D^p) =$	
$D_6 = \{s_1\}$	0.16		
$D_7 = \{s_2\}$	0.04		
$D_8 = \emptyset$	0.04		

answer	prob.
$\{p'\}$	0.54
$\emptyset$	0.46

(a)

(c)

Figure 3: (a) The possible worlds for D^p in Figure 2, (b) a query q , and (c) its possible answers.

them: using heuristics to chose a plan that avoids large errors, and using a Monte-Carlo simulation algorithm, which is more expensive but can guarantee arbitrarily small errors.

Outline We give motivating examples in Sec. 2, define the problem in Sec. 3, and describe our techniques in Sec. 4-8. Sec. 9 reports experiments and Sec. 10 describes related work. We conclude in Sec. 11.

2 Examples

We illustrate the main concepts and techniques of this paper with two simple examples.

Probabilistic Database In a probabilistic database each tuple has a certain probability of belonging to the database. Figure 2 shows a probabilistic database D^p with two tables, S^p and T^p : the tuples in S^p have probabilities 0.8 and 0.5, and the unique tuple in T^p has probability 0.6. We use the superscript p to emphasize that a table or a database is probabilistic. We assume in this example that the tuples are independent probabilistic events, in which case the database is called *extensional* [11].

The meaning of a probabilistic database is a probability distribution on all database instances, which we call *possible worlds*, and denote $pwd(D^p)$. Fig. 3 (a) shows the eight possible instances with non-zero probabilities, which are

computed by simply multiplying the tuple probabilities, as we have assumed them to be independent. For example, the probability of D_2 is $0.8 * (1 - 0.5) * 0.6 = 0.24$, since the instance contains the tuples s_1 and t_1 and does not contain s_2 .

We now illustrate query evaluation on probabilistic databases. Consider the conjunctive query q in Fig. 3 (b). Its meaning on D^p is a set of possible answers, shown in Fig. 3 (c). It is obtained by applying q to each deterministic database in $pwd(D^p)$, and adding the probabilities of all instances that return the same answer. In our example we have $q(D_1) = q(D_2) = q(D_3) = \{p'\}$, and $q(D_4) = \dots = q(D_8) = \emptyset$. Thus, the probability of the answer being $\{p'\}$ is $0.24 + 0.24 + 0.06 = 0.54$, while that of the answer $\emptyset$ is 0.46. This defines the set of possible answers, denoted $q^{pwd}(D^p)$. Notice that we have never used the structure of the query explicitly, but only applied it to deterministic databases taken from $pwd(D^p)$. Thus, one can give a similar semantics to *any* query q , no matter how complex, because we only need to know its meaning on deterministic databases.

The set of possible answers $q^{pwd}(D^p)$ may be very large, and it is impractical to return it to the user. Instead, we compute for each possible tuple t a *probability rank* that t belongs to any answer, and sort tuples sorted by this rank. We denote this $q^{rank}(D^p)$. In our example this is:

$$q^{rank}(D^p) = \begin{array}{|c|c|} \hline \mathbf{D} & \mathbf{Rank} \\ \hline 'p' & 0.54 \\ \hline \end{array}$$

In this simple example $q^{rank}(D^p)$ contains a single tuple and the distinction between q^{pwd} and q^{rank} is blurred. To see this distinction clearer, consider another query, $q_1(x) : -S^p(x, y), T^p(z, y), y = z$. Here q_1^{pwd} and q_1^{rank} are given by:

$$q_1^{pwd}(D^p) = \begin{array}{|c|c|} \hline \text{answer} & \text{probability} \\ \hline \{'m', 'n'\} & 0.24 \\ \{'m'\} & 0.24 \\ \{'n'\} & 0.06 \\ \emptyset & 0.46 \\ \hline \end{array} \quad q_1^{rank}(D^p) = \begin{array}{|c|c|} \hline \mathbf{D} & \mathbf{Rank} \\ \hline 'm' & 0.48 \\ 'n' & 0.30 \\ \hline \end{array}$$

For example, the rank probability of $'m'$ is obtained as $Pr(\{'m', 'n'\}) + Pr(\{'m'\})$. In general, $q^{pwd}(D^p)$ may be exponentially large, while $q^{rank}(D^p)$ is simply a set of tuples, which are sorted by **Rank**. The problem in this paper is now to compute $q^{rank}(D^p)$ efficiently.

Extensional Query Semantics A natural attempt to compute $q^{rank}(D^p)$ is to represent q as a query plan then compute the probabilities of all tuples in all intermediate results. For the query q in Fig. 3 (b), such a plan is $p = \Pi_{\mathbf{D}}(S^p \bowtie_{\mathbf{B}=\mathbf{C}} T^p)$, and the corresponding probabilities are shown in Fig. 4. The formulas for the probabilities assume tuple independence, are taken from [11] and are rather straightforward (we review them in Sec. 4). For example the probability of a joined tuple $s \bowtie t$ is the product of the probabilities of s and

A	B	C	D
'm'	1	1	'p'
'n'	1	1	'p'

$$\begin{array}{l}
\text{prob} \\
0.8 * 0.6 = 0.48 \\
0.5 * 0.6 = 0.30
\end{array}$$

(a) $S^p \bowtie_{B=C} T^p$

D
'p'

$$\begin{array}{l}
\text{prob} \\
(1 - (1 - 0.48)(1 - 0.3)) = 0.636
\end{array}$$

(b) $\Pi_D(S^p \bowtie_{B=C} T^p)$

Figure 4: Evaluation of $\Pi_D(S^p \bowtie_{B=C} T^p)$

B
1

$$\begin{array}{l}
\text{prob} \\
(1 - (1 - 0.8)(1 - 0.5)) = 0.9
\end{array}$$

(a) $\Pi_B(S^p)$

B	C	D
1	1	'p'

$$\begin{array}{l}
\text{prob} \\
0.9 * 0.6 = 0.54
\end{array}$$

(b) $\Pi_B(S^p) \bowtie_{B=C} T^p$

D
'p'

$$\begin{array}{l}
\text{prob} \\
0.54
\end{array}$$

(c) $\Pi_D(\Pi_B(S^p) \bowtie_{B=C} T^p)$

Figure 5: Evaluation of $\Pi_D(\Pi_B(S^p) \bowtie_{B=C} T^p)$

t. Clearly, this approach is much more efficient than computing the possible worlds $q^{pvd}(D^p)$ and then computing $q^{rank}(D^p)$, but it is wrong ! It's answer is 0.636, while it should be 0.54. The reason is that the two tuples in $S^p \bowtie_{B=C} T^p$ are not independent events, hence the formula used in Π_D is wrong.

However, let us consider an alternative plan, $p' = \Pi_D((\Pi_B(S^p)) \bowtie_{B=D} T^p)$. The extensional evaluation of this expression is shown in Figure 5, and this time we do get the correct answer. As we will show later, this plan will always compute the correct answer to q , on any probabilistic tables S^p, T^p . In this paper we show how to find automatically a plan whose extensional evaluation returns the correct answer to a query q . Finding such a plan requires pushing projections early (as shown in this example), join reordering, and other kinds of rewritings.

Queries with uncertain matches While query evaluation on probabilistic

databases is an important problem in itself, our motivation comes from answering SQL queries with uncertain matches, and ranking their results. We illustrate here with a simple example on the Stanford movie database[1].

```

SELECT  DISTINCT F.title, F.year
FROM    Director D, Films F
WHERE   D.did = F.did
          and    D.name  $\approx$  'Coppola'
          and    F.title  $\approx$  'rain man'
          and    F.year  $\approx$  1995

```

The predicates on the director name and the movie title and year are here *uncertain*.

Our approach is to translate the query into a regular query over a probabilistic databases. Each tuple in the table **Films** is assigned a probability based on how well it matches the predicates **title** $\approx$ 'rain man' and **year** $\approx$ 1995. Several techniques for doing this exist already, and in this paper we will adopt existing ones: see Sec. 8. In all cases, the result is a probabilistic table, denoted **Films**^{*p*}. Similarly, the uncertain predicate on **Director** generates a probabilistic table **Director**^{*p*}. Then, we evaluate the following query:

```

SELECT  DISTINCT F.title, F.year
FROM    Directorp D, Filmsp F
WHERE   D.did = F.did

```

This is similar to the query *q* considered earlier (Figure 3 (b)), and the same extensional plan can be used to evaluate it. Our system returns:

title	year	rank
The Rainmaker (by Coppola)	1997	0.110
The Rain People (by Coppola)	1969	0.089
Rain Man (by Levinson)	1988	0.077
Finian's Rainbow (by Coppola)	1968	0.069
Tucker, Man and Dream (Coppola)	1989	0.061
Rain or Shine (by Capra)	1931	0.059
...	...	...

3 Problem Definition

We review here the basic definitions in probabilistic databases, based on [11, 4], and state our problem.

Basic Notations We write *R* for a relation name, *Attr*(*R*) for its attributes, and $r \subseteq U^k$ for a relation instance where *k* is *arity*(*R*) and *U* is a fixed, finite universe. $\bar{R} = R_1, \dots, R_n$ is a database schema and *D* denotes a database instance. We write $\Gamma \models D$ when *D* satisfies the functional dependencies in Γ .

Probabilistic Events Let *AE* be a set of symbols and $Pr : AE \rightarrow [0, 1]$ a probability function. Each element of *AE* is called a *basic event*, and we assume that all basic events are independent. The event $\perp \in AE$ denotes the impossible event and $Pr(\perp) = 0$. A *complex event* is an expression constructed from atomic

events using the operators $\wedge, \vee, \neg$. E denotes the set of all complex events. For each complex event e , let $Pr(e)$ be its probability.

Example 3.1 Consider $e = (s_1 \wedge t_1) \vee (s_2 \wedge t_1)$, and assume $Pr(s_1) = 0.8$, $Pr(s_2) = 0.5$, $Pr(t_1) = 0.6$. To compute $Pr(e)$ we construct the truth table for $e(s_1, s_2, t_1)$ and identify the entries where e is true, namely $(1, 0, 1), (0, 1, 1), (1, 1, 1)$. The three entries have probabilities $Pr(s_1)(1 - Pr(s_2))Pr(t_1) = 0.8 \times 0.5 \times 0.6 = 0.24$, $(1 - Pr(s_1))Pr(s_2)Pr(t_1) = 0.06$ and $Pr(s_1)Pr(s_2)Pr(t_1) = 0.24$ respectively. Then $Pr(e)$ is their sum, 0.54.

This method generalizes to any complex event $e(s_1, \dots, s_k)$, but it is important to notice that this algorithm is exponential in k . This cannot be avoided: it is known that computing $Pr(e)$ is #P-complete [24] even for complex events without negation.

Probabilistic Databases A *probabilistic relation* is a relation with a distinguished *event attribute* E , whose value is a complex event. We add the superscript p to mean “probabilistic”, i.e. write $R^p, r^p, \bar{R}^p, \Gamma^p$. Given R^p , we write R for its “deterministic” part, obtained by removing the event attribute: $Attr(R) = Attr(R^p) - \{E\}$. Users “see” only R , but the system needs to access the event attribute $R^p.E$. The set of functional dependencies Γ^p always contains:

$$Attr(R) \rightarrow R^p.E$$

for every relation R^p : this ensures that we don’t associate two different events e_1 and e_2 to the same tuple t (instead, we may want to associate $e_1 \vee e_2$ to t).

In addition to this *tabular* representation of a probabilistic relation, we consider a *functional* representation, where a probabilistic instance r^p , of type R^p , is described by the following function $e_R : U^k \rightarrow E$, where $k = arity(R)$. When t occurs in r^p together with some event e , then $e_R(t) = e$, otherwise $e_R(t) = \perp$. Conversely, one can recover r^p from the function e_R by collecting all tuples for which $e_R(t) \neq \perp$.

The input probabilistic databases we consider have only atomic events: complex events are introduced only by query evaluation. A probabilistic relation with atomic events which satisfies the FD $R^p.E \rightarrow Attr(R)$ is called *extensional*. Otherwise, it is called *intensional*. For example, the database in Fig. 2 is an extensional probabilistic database, where the atomic events are s_1, s_2, t_1 respectively.

Semantics of a probabilistic database We give a simple and intuitive meaning to a probabilistic relation based on possible worlds. The meaning of a probabilistic relation r^p of type R^p is a probability distribution on deterministic relations r of type R , which we call the *possible worlds*, and denote $pwd(r^p)$. Let $e_R : U^k \rightarrow E$ be the functional representation of r^p . Given $r \subseteq U^k$, $Pr(r)$ is defined to be $Pr(\bigwedge_{t \in r} e_R(t)) \wedge (\bigwedge_{t \notin r} \neg e_R(t))$. Intuitively, this is the probability that exactly the tuples in r are “in” and all the others are “out”. One can check that $\sum_{r \subseteq U^k} Pr(r) = 1$. If we have a set of functional dependencies Γ^p , denote Γ its projection to the deterministic attributes. If $\Gamma^p \models r^p$, we want $Pr(r)$ to

be a probability distribution on those instances which satisfy r : this is the case, because if $\Gamma \not\models r$, then $Pr(r) = 0$, hence $\sum_{r \subseteq U^k, \Gamma \models r} Pr(r) = 1$.

Similarly, the meaning of a probabilistic database D^p is a probability distribution on all deterministic databases D , denoted $pwd(D^p)$.

Query semantics Let q be a query of arity k over a deterministic schema $\bar{R}$. We define a very simple and intuitive semantics for the query. Users think of q as normal query on a deterministic database, but the database is given by a probability distribution rather than being fixed. As a result, the query's answer is also a probability distribution. Formally, given a query q and a probabilistic database D^p : $q^{pwd}(D^p)$ is the following probability distribution on all possible answers, $Pr_q : \mathcal{P}(U^k) \rightarrow [0, 1]$:

$$\forall S \subseteq U^k, Pr_q(S) = \sum_{D|q(D)=S} Pr(D)$$

We call this the *possible worlds semantics*. This definition makes sense for every query q that has a well defined semantics on all deterministic databases.

It is impossible to return $q^{pwd}(D^p)$ to the user. Instead, we compute a *probabilistic ranking* on all tuples $t \in U^k$, defined by the function: $rank_q(t) = \sum_S \{Pr_q(S) \mid S \subseteq U^k, t \in S\}$, for every tuple $t \in U^k$. We denote with $q^{rank}(D^p)$ a tabular representation of the function $rank_q$: this is a table with $k + 1$ attributes, where the first k represent a tuple in the query's answer while the last attribute, called **Rank** is a real number in $[0, 1]$ representing its probability.

The Query Evaluation Problem This paper addresses the following problem: given schema $\bar{R}^p, \Gamma^p$, a probabilistic database D^p and a query q over schema $\bar{R}$, compute the probabilistic rankings $q^{rank}(D^p)$.

Application to queries with uncertain predicates Consider now a deterministic database D and a query $q^\sim$ that explicitly mentions some uncertain predicates. We convert this problem into evaluating a deterministic query q , obtained by removing all uncertain predicates from $q^\sim$, on a probabilistic database, obtained by associating a probability $Pr(t)$ to each tuple t based on how well t satisfies the uncertain predicates in the query.

4 Query Evaluation

We turn now to the central problem, evaluating $q^{rank}(D^p)$ for a query q , and a probabilistic database D^p . Applying the definition directly is infeasible, since it involves iterating over a large set of database instances. Instead, we will first review the intensional evaluation of [11] then describe our approach.

We restrict our discussion first to conjunctive queries, or, equivalently select (distinct)-project-join queries. This helps us better understand the query evaluation problem and its complexity, and will consider more complex query expressions in Sec. 7. We use either datalog notation for our queries q , or plans p in the select/project/product algebra³: $\sigma, \Pi, \times$.

³Notice that Π also does duplicate elimination

4.1 Intensional Query Evaluation

One method for evaluating queries on probabilistic databases is to use complex events, and was introduced in [11]. We review it here and discuss its limitations. Start by expressing q as a query plan, using the operators $\sigma, \Pi, \times$. Then modify each operator to compute the event attribute E in each intermediate result: denote $\sigma^i, \Pi^i, \times^i$ the modified operators. It is more convenient to introduce them in the functional representation, by defining the complex event $e_p(t)$ for each tuple t , inductively on the query plan p :

$$\begin{aligned} e_{\sigma_c^i(p)}(t) &= \begin{cases} e_p(t) & \text{if } c(t) \text{ is true} \\ \perp & \text{if } c(t) \text{ is false} \end{cases} \\ e_{\Pi_{\bar{A}}^i(p)}(t) &= \bigvee_{t': \Pi_{\bar{A}}(t')=t} e_p(t') \\ e_{p \times^i p'}(t, t') &= e_p(t) \wedge e_{p'}(t') \end{aligned} \tag{1}$$

The tabular definitions for $\sigma^i, \Pi^i, \times^i$ follow easily: σ^i acts like σ then copies the complex events from the input tuples to the output tuples; Π^i associates to a tuple t the complex event $e_1 \vee \dots \vee e_n$ obtained from the complex events of all input tuples $t_1, \dots, t_n$ that project into t ; and $\times^i$ simply associates to a product tuple (t, t') the complex event $e \wedge e'$.

Example 4.1 Let us consider the database D^p described in Figure 2. Consider the query plan, $p = \Pi_D(S^p \bowtie_{B=C} T^p)$. Figure 6 shows the intensional evaluation of the query (we used the tuple names as atomic events). $p^i(D^p)$ contains a single tuple $'p'$ with the event $(s_1 \wedge t_1) \vee (s_2 \wedge t_1)$.

Thus, $p^i(D^p)$ denotes an intensional probabilistic relation. It can be shown that this is independent on the particular choice of plan p , and we denote $q^i(D^p)$ the value $p^i(D^p)$ for any plan p for q , and call it the *intensional semantics*⁴ of q on the probabilistic database D^p . We prove now that it is equivalent to the possible worlds semantics, $q^{pwd}(D^p)$.

Theorem 4.2. *The intensional semantics and the possible worlds semantics on probabilistic databases are equivalent for conjunctive queries. More precisely, $pwd(q^i(D^p)) = q^{pwd}(D^p)$ for every intensional probabilistic database D^p and every conjunctive query q .*

Proof. Every tuple t in $q^i(D^p)$ has a complex event $t.E$ associated with it. $pwd(q^i(D^p))$ consists of a set of worlds, each world assigning a truth value to the set of atomic events. A tuple t belongs to a world in $pwd(q^i(D^p))$ if $t.E$ is true in that world.

$q^{pwd}(D^p)$ also consists of a set of worlds, each assigning a truth value to the set of atomic events. The content of a world in $q^{pwd}(D^p)$ is the output of q on the database defined by that world.

⁴In [11] this is the only query semantics considered.

A	B	C	D	E
'm'	1	1	'p'	$s_1 \wedge t_1$
'n'	1	1	'p'	$s_2 \wedge t_1$

(a) $S^p \bowtie_{B=C}^i T^p$

D	E
'p'	$(s_1 \wedge t_1) \vee (s_2 \wedge t_1)$

(b) $\Pi_D^i(S^p \bowtie_{B=C}^i T^p)$

D	Rank
'p'	$Pr((s_1 \wedge t_1) \vee (s_2 \wedge t_1)) = 0.54$

(c) $q^{rank}(D^p) = Pr(\Pi_D^i(S^p \bowtie_{B=C}^i T^p))$

Figure 6: Intensional Evaluation of $\Pi_D(S^p \bowtie_{B=C} T^p)$

Given a world W defined by an assignment of truth values to the atomic variables, let $q^{pwd}(D^p)[W]$ and $pwd(q^i(D^p))[W]$ denote the contents of the corresponding worlds.

We will prove, by induction on the size of q , that for all W , $q^{pwd}(D^p)[W] = pwd(q^i(D^p))[W]$.

As a base case, if a query just selects a single relation, it holds trivially. If q is a larger query, there are three possibilities:

1. $q = \Pi_A(q_1)$. Consider any world W . By induction hypothesis,

$$pwd(q_1^i(D^p))[W] = q_1^{pwd}(D^p)[W]$$

Given a tuple t in $\Pi_A(q_1)$, its event $t.E$ is given by $\vee_{t_1: \Pi_A(t_1)=t} t_1.E$.

$$\begin{aligned}
t \in q^{pwd}(D^p)[W] &\Leftrightarrow t.E = \text{true} \\
&\Leftrightarrow (\vee_{t_1: \Pi_A(t_1)=t} t_1.E) = \text{true} \\
&\Leftrightarrow \exists t_1, \Pi_A(t_1) = t, \\
&\quad t_1.E \in pwd(q_1^i(D^p))[W] \\
&\Leftrightarrow \exists t_1, \Pi_A(t_1) = t, \\
&\quad t_1.E \in q_1^{pwd}(D^p)[W] \\
&\Leftrightarrow t \in q^{pwd}(D^p)[W]
\end{aligned}$$

2. $q = \sigma_c(q_1)$. Consider any world W . Again, we have

$$pwd(q_1^i(D^p))[W] = q_1^{pwd}(D^p)[W]$$

t belongs to $pwd(q^i(D^p))[W]$ iff $t.E$ is true, i.e. t satisfies σ_c and t belongs to $pwd(q_1^i(D^p))[W]$. Similarly, t belongs to $pwd(q^i(D^p))[W]$ iff t satisfies σ_c and t belongs to $q_1^{pwd}(D^p)[W]$. Therefore, we get $pwd(q^i(D^p))[W] = q^{pwd}(D^p)[W]$.

3. $q = q_1 \bowtie q_2$. We have

$$\begin{aligned} pwd(q_1^i(D^p))[W] &= q_1^{pwd}(D^p)[W], \\ pwd(q_2^i(D^p))[W] &= q_2^{pwd}(D^p)[W] \end{aligned}$$

Given a tuple $t = (t_1, t_2)$ belonging to q , $t.E = t_1.E \wedge t_2.E1$. Thus,

$$\begin{aligned} t \in q^{pwd}(D^p)[W] &\Leftrightarrow t.E = \text{true} \\ &\Leftrightarrow t_1.E = \text{true}, t_2.E = \text{true} \\ &\Leftrightarrow t_1 \in pwd(q_1^i(D^p))[W], \\ &\quad t_2 \in pwd(q_2^i(D^p))[W] \\ &\Leftrightarrow t_1 \in q_1^{pwd}(D^p)[W], \\ &\quad t_2 \in q_2^{pwd}(D^p)[W] \\ &\Leftrightarrow t \in q^{pwd}(D^p)[W] \end{aligned}$$

Thus, by induction, $pwd(q^i(D^p))$ and $q^{pwd}(D^p)$ are equal. $\square$

Theorem 4.2 allows us to compute $q^{rank}(D^p)$, as follows. First compute $q^i(D^p)$, then compute the probability $Pr(e)$ for each complex event. Then $q^{rank}(D^p) = Pr(q^i(D^p))$.

Example 4.3 Fig. 6(c) shows $q^{rank}(D^p)$ for Ex. 4.1. $Pr((s_1 \wedge t_1) \vee (s_2 \wedge t_1))$ was shown in Ex. 3.1.

It is very impractical to use the intensional semantics to compute the rank probabilities, for two reasons. First, the event expressions in $q^i(D^p)$ can become very large, due to the projections. In the worst case the size of such an expression can become of the same order of magnitude as the database. This increases the complexity of the query operators significantly, and makes the task of an optimizer much harder, because now the cost per tuple is not longer constant. Second, for each tuple t one has to compute $Pr(e)$ for its event e , which is a #P-complete problem.

4.2 Extensional Query Evaluation

We now modify the query operators to compute probabilities rather than complex events: we denote $\sigma^e, \Pi^e, \times^e$ the modified operators. This is much more efficient, since it involves manipulating real numbers rather than event expressions. We define a number $Pr_p(t) \in [0, 1]$ for each tuple t , by induction on

the structure of the query plan p . The inductive definitions below should be compared with those in Equations (1).

$$\begin{aligned}
Pr_{\sigma^e(p)}(t) &= \begin{cases} Pr_p(t) & \text{if } c(t) \text{ is true} \\ 0 & \text{if } c(t) \text{ is false} \end{cases} \\
Pr_{\Pi_A^e(p)}(t) &= 1 - \prod_{t': \Pi_A(t')=t} (1 - Pr_p(t')) \\
Pr_{p \times^e p'}(t, t') &= Pr_p(t) \times Pr_{p'}(t')
\end{aligned}$$

Again, the tabular definitions of $\sigma^e, \Pi^e, \times^e$ follow easily: σ^e acts like σ then propagates the tuples' probabilities from the input to the output, Π^e computes the probability of a tuples t as $1 - (1 - p_1)(1 - p_2) \dots (1 - p_n)$ where $p_1, \dots, p_n$ are the probabilities of all input tuples that project to t , while $\times$ computes the probability of each tuple (t, t') as $p \times p'$.

Thus, $p^e(D^p)$ is an extensional probabilistic relation, which we call the *extensional semantics* of the plan p . If we know $p^e(D^p) = q^{rank}(D^p)$, then we simply execute the plan under the extensional semantics. But, unfortunately, this is not always the case, as we saw in Sec. 2. Moreover, $p^e(D^p)$ depends on the particular plan p chosen for q . Our goal is to find a plan for which the extensional semantics is correct.

Definition 4.4. *Given a schema $\bar{R}^p, \Gamma^p$, a plan p for a query q is safe if $p^e(D^p) = q^{rank}(D^p)$ for all D^p of that schema.*

We show next how to find a safe plan.

4.3 The Safe-Plan Optimization Algorithm

We use the following notations for conjunctive queries:

- $Rel(q) = \{R_1, \dots, R_k\}$ all relation names occurring in q . We assume that each relation name occurs at most once in the query (more on this in Sec. 7).
- $PRel(q)$ = the probabilistic relation names in q , $PRel(q) \subseteq Rel(q)$.
- $Attr(q)$ = all attributes in all relations in q . To disambiguate, we denote attributes as $R_i.A$.
- $Head(q)$ = the head attributes in q , $Head(q) \subseteq Attr(q)$.

Let q be a conjunctive query. We define the *induced* functional dependencies $\Gamma^p(q)$ on $Attr(q)$:

- Every FD in Γ^p is also in $\Gamma^p(q)$.
- For every join predicate $R_i.A = R_j.B$, both $R_i.A \rightarrow R_j.B$ and $R_j.B \rightarrow R_i.A$ are in $\Gamma^p(q)$.

- For every selection predicate $R_i.A = c$, $\emptyset \rightarrow R_i.A$ is in $\Gamma^p(q)$.

We seek a safe plan p , i.e. one that computes the probabilities correctly. For that each operator in p must be safe, i.e. compute correct probabilities: we define this formally next.

Let q_1, q_2 be two queries, and let $op \in \{\sigma, \Pi, \times\}$ be a relational operator. Consider the new query $op(q_1, q_2)$ (or just $op(q_1)$ when op is unary). We say that op^e is *safe* if $op^e(Pr(q_1^i(D^p)), Pr(q_2^i(D^p))) = Pr(op^i(q_1^i(D^p)), q_2^i(D^p))$ (and similarly for unary operators), $\forall D^p$ s.t. $\Gamma^p \models D^p$. In other words, op is safe if, when given correct probabilities for its inputs op^e computes correct probabilities for the output tuples.

Theorem 4.5. *Let q, q' be conjunctive queries.*

1. σ_c^e is always safe in $\sigma_c(q)$.
2. $\times^e$ is always safe in $q \times q'$.
3. $\Pi_{A_1, \dots, A_k}^e$ is safe in $\Pi_{A_1, \dots, A_k}(q)$ iff for every $R^p \in PRel(q)$ the following can be inferred from $\Gamma^p(q)$:

$$A_1, \dots, A_k, R^p.E \rightarrow Head(q) \quad (2)$$

Proof.

1. Follows trivially from definition.
2. Since we assume all relations in the query to be distinct, the complex events in the output of q and q' comprise of distinct atomic events. Thus, given a tuple $t_{join} = (t, t') \in q \times q'$, $Pr(t_{join}.E) = Pr(t.E \wedge t'.E) = Pr(t.E)Pr(t'.E)$. So, the join operator is safe.
3. For each output tuple t of the project, consider the set S_t of input tuples that map to t . The operator is safe if and only if, for each such t , the complex events corresponding to tuples in S_t are independent. Thus, among all tuples having the same value for $A_1, \dots, A_k$, no basic event (i.e. $R^p.E$ for some probabilistic relation R^p) occurs in two of them having different values of $Head(q)$. Thus, the following functional dependency must hold for each R^p .

$$A_1, \dots, A_k, R^p.E \rightarrow Head(q)$$

□

Obviously, a plan p consisting of only safe operators is safe. As we prove below, the converse also holds.

Theorem 4.6. *Let p be a safe relational algebra plan for a query q consisting of selects, projects and joins. Then, all operators in p are safe.*

Proof. We will prove the following statement: If p is a plan for a query q that has atleast one unsafe operator and t is any tuple of the same arity as $Heads(q)$, there is a database D^p such that $p^e(D^p)$ consists of a single tuple t with incorrect probability.

First, it is easy to see that for any plan p and database D^p , $p^e(D^p) \geq q^{rank}(D^p)$. This is because the only operators that are unsafe are projects, and they can only overestimate the probabilities. Thus, there can only be one-sided errors in the probabilities.

We will prove the theorem by induction on the size of the plan.

If p returns a single relation, its safe and the assertion holds trivially.

Consider a larger plan p and assume it is safe. There are three cases:

1. $p = q_1 \bowtie q_2$. Since p is unsafe, at least one of q_1 and q_2 is unsafe. W.L.O.G, assume q_1 is unsafe. Given any tuple $t = (t_1, t_2)$, a database can be constructed such that q_1 results in a single tuple t with incorrect probability. Since $Pr_p(t_1, t_2) = Pr_{p_1}(t_1) * Pr_{p_2}(t_2)$, for the final probability to be correct, both $Pr_{p_1}(t_1)$ and $Pr_{p_2}(t_2)$ have to be correct (since errors can only be one-sided). Thus, we have constructed a database that results in a single tuple t with incorrect probability.
2. $p = \Pi_A(q_1)$. If q_1 is safe, the final project operator must be unsafe and hence, there exists a database where the final result is incorrect. If t is any tuple in the output whose probability is incorrect, we can restrict the database to produce only the tuple t (whose probability is still incorrect).
If q_1 is not safe, there is a database on which q_1 produces a single tuple with wrong probability. Thus, the final project also produces a single tuple whose probability is still wrong.
3. $p = \sigma_c(q_1)$. Consider any tuple t that satisfies the condition c and create a database on which q_1 produces t with incorrect probability. Thus, $\sigma_c(q_1)$ also produces a single tuple with incorrect probability.

□

We explain safe plans with an example below.

Example 4.7 Continuing the example in Sec. 2, assume that both S^p and T^p are extensional probabilistic relations, hence Γ^p is:

$$\begin{aligned}
S^p.A, S^p.B &\rightarrow S^p.E \\
T^p.C, T^p.D &\rightarrow T^p.E \\
S^p.E &\rightarrow S^p.A, S^p.B \\
T^p.E &\rightarrow T^p.C, T^p.D
\end{aligned}$$

The last two dependencies hold because the relations are extensional. Consider the plan $\Pi_D(S^p \bowtie_{B=C} T^p)$. We have shown in Fig. 4 that, when evaluated extensionally, this plan is incorrect. We explain here the reason: the operator

Π_D^e is not safe. An intuitive justification can be seen immediately by inspecting the intensional relation $S^p \bowtie_{B=C}^i T^p$ in Fig. 6 (a). The two complex events share the common atomic event t_1 , hence they are correlated probabilistic events. But the formula for Π_D^e only works when these events are independent. We show how to detect formally that Π_D^e is unsafe. We need to check:

$$\begin{aligned} T^p.D, S^p.E &\rightarrow S^p.A, S^p.B, T^p.C, T^p.D \\ T^p.D, T^p.E &\rightarrow S^p.A, S^p.B, T^p.C, T^p.D \end{aligned}$$

The first follows from Γ^p and from the join condition $B = C$, which adds $S^p.B \rightarrow T^p.C$ and $T^p.C \rightarrow S^p.B$. But the second fails: $T^p.D, T^p.E \not\rightarrow S^p.A$.

Example 4.8 Continuing the example, consider now the plan $\Pi_D(\Pi_B(S^p) \bowtie_{B=C} T^p)$. We will prove that Π_D^e is safe. For that we have to check:

$$\begin{aligned} T^p.D, S^p.E &\rightarrow S^p.B, T^p.C, T^p.D \\ T^p.D, T^p.E &\rightarrow S^p.B, T^p.C, T^p.D \end{aligned}$$

Both hold, hence Π_D^e is safe. Similarly, Π_B^e is safe in $\Pi_B(S^p)$, which means that the entire plan is safe.

Before we describe our algorithm for finding a safe plan, we need some terminology.

Definition 4.9. (*Separate relations*) For a query φ , two relations R_i and R_j are *separate* if none of the columns of one table that are used in the join with the other table are projected out in φ .

Definition 4.10. (*Separation*) Two sets of relations S_i and S_j are said to form a separation for query φ iff

1. They partition the set $PRels(\varphi)$
2. For any pair of relations R_i and R_j such that $R_i \in S_i$ and $R_j \in S_j$, they must be separate.

Algorithm 1 is our optimization algorithm for finding a safe plan. It proceeds top-down, as follows. First, it tries to do all safe projections late in the query plan. When no more late safe projections are possible for a query q , then it tries to perform a join $\bowtie_c$ instead, by splitting q into $q_1 \bowtie_c q_2$. Since $\bowtie_c$ is the last operation in the query plan, all attributes in c must be in $Head(q)$. Hence, $Rels(q_1)$ and $Rels(q_2)$ must form a separation.

Given a query q , a separation can be found as follows. Construct a graph G whose nodes are $Rels(q)$ and whose edges are all pairs (R_i, R_j) s.t. q contains

some join condition $R_i.A = R_j.B$ with both⁵ $R_i.A$ and $R_j.B$ in $Head(q)$. Find the connected components of G , and choose q_1 and q_2 to be any partition of these connected components: this defines $Rel_s(q_i)$ and $Attr(q_i)$ for $i = 1, 2$. Define $Head(q_i) = Head(q) \cap Attr(q_i)$, for $i = 1, 2$. If G is a connected graph, then the query has no safe plans (more on this below). If G has multiple connected components, then we have several choices for splitting q , and we can deploy any standard cost based optimizations algorithm that works in top-down fashion⁶.

Finally, the algorithm terminates when no more projections are needed. The remaining join and/or selection operators can be done in any order.

Algorithm 1 SAFE-PLAN(q)

```

1: if  $Head(q) = Attr(q)$  then
2:   return any plan  $p$  for  $q$ 
3:   ( $p$  is projection-free, hence safe)
4: end if
5: for  $A \in (Attr(q) - Head(q))$  do
6:   let  $q_A$  be the query obtained from  $q$ 
7:   by adding  $A$  to the head variables
8:   if  $\Pi_{Head(q)}(q_A)$  is a safe operator then
9:     return  $\Pi_{Head(q)}(\text{SAFE-PLAN}(q_A))$ 
10:  end if
11: end for
12: Split  $q$  into  $q_1 \bowtie_c q_2$  (see text)
13: if no such split exists then
14:   return error("No safe plans exist")
15: end if
16: return  $\text{SAFE-PLAN}(q_1) \bowtie_c \text{SAFE-PLAN}(q_2)$ 

```

Example 4.11 Continuing the example in Sec. 2, consider the original query in Fig. 3 (b), which we rewrite now as:

$$q(D) : -S^p(A, B), T^p(C, D), B = C$$

Here $Attr(q) = \{A, B, C, D\}$ and $Head(q) = \{D\}$ (we write D instead of $T^p.D$, etc, since all attributes are distinct). The algorithm first considers the three attributes A, B, C in $Attr(q) - Head(q)$, trying to see if they can be projected out late in the plan. A cannot be projected out. Indeed, the corresponding q_A is:

$$q_A(A, D) : -S^p(A, B), T^p(C, D), B = C$$

and Π_D^e is unsafe in $\Pi_D(q_A)$ because $T^p.D, T^p.E \not\rightarrow S^p.A$, as we saw in Example 4.7. However, the other two attributes can be projected out, hence the plan

⁵One can show that, if $R_i.A$ is in $Head(q)$, then so is $R_j.B$. Indeed, assume $R_j.B \notin Head(q)$. Then $\Pi_{Head(q)}(q_{R_j.B})$ is safe, so we should have performed it first. Then, both $R_i.A$ and $R_j.B$ are in $Head(q_{R_j.B})$.

⁶It is also possible to adapt our algorithm to work with a bottom-up optimizer.

for q is $\Pi_D(q_{BC})$, where:

$$q_{BC}(B, C, D) : -S^p(A, B), T^p(C, D), B = C$$

Now we optimize q_{BC} , where $Attr(q_{BC}) = \{A, B, C, D\}$, $Head(q_{BC}) = \{B, C, D\}$. No projection is possible, but we can split the query into $q_1 \bowtie_{B=C} q_2$ where q_1, q_2 are:

$$\begin{aligned} q_1(B) &: - S^p(A, B) \\ q_2(C, D) &: - T^p(C, D) \end{aligned}$$

The split $q_{BC} = q_1 \bowtie_{B=C} q_2$ is indeed possible since both B and C belong to $Head(q_{BC})$. Continuing with q_1, q_2 , we are done in q_2 , while in q_1 we still need to project out A , $q_1 = \Pi_B(S^p)$, which is safe since $B, S^p.E \rightarrow A$. Putting everything together gives us the following safe plan: $p' = \Pi_D(\Pi_B(S^p) \bowtie_{B=C} T^p)$.

We state now the soundness of our algorithm: the proof follows easily from the fact that all projection operators are safe. We prove in the next section that the algorithm is also complete.

Proposition 4.12. *The SAFE-PLAN optimization algorithm is sound, i.e. any plan it returns is safe.*

Proof. Define the size of the query q to be $|Rels(q)| + |Attrs(q)| - |Head(q)|$. We will prove the proposition by induction on the size of the query.

Algorithm 1 returns a plan in the following three cases:

1. It returns at line 2. In this case, the plan consists only of joins and selects, and hence, is safe.
2. It returns at line 9. Note that q_A is smaller in size than q . By induction, $Safe - Plan(q_A)$ is safe. Also, by definition, the final project operator is safe. Thus, the algorithm returns a safe plan.
3. The algorithm returns at line 16. q_1 and q_2 are both smaller than q . By induction, both $Safe - Plan(q_1)$ and $Safe - Plan(q_2)$ are safe plans. These plans are then connected by a join operator, which is always safe. So the returned plan is safe.

□

4.4 Completeness of Safe-Plan algorithm

We have shown that SAFE-PLAN algorithm is sound. We next prove that it is complete, i.e., if the query is safe, the algorithm finds a safe plan for it.

We start with few basic results about extensional evaluation.

Lemma 4.13. *Under extensional evaluation, for any query q ,*

$$\Pi_A^e(\Pi_{A \cup B}^e(q)) = \Pi_A^e(q)$$

Proof. These two expressions have the same output tuples. We only have to show that the tuples have the same probability in both expressions.

Consider a tuple t belonging to $\Pi_A(\Pi_{A \cup B}(q))$. Then,

$$\begin{aligned}
& 1 - Pr_{\Pi_A(\Pi_{A \cup B}(q))}(t) \\
= & \prod_{(t' | \Pi_A(t')=t)} (1 - Pr_{\Pi_{A \cup B}(q)}(t')) \\
= & \prod_{(t' | \Pi_A(t')=t)} \prod_{(t'' | \Pi_{A \cup B}(t'')=t')} (1 - Pr_q(t'')) \\
= & \prod_{(t'' | \Pi_A(t'')=t)} (1 - Pr_q(t'')) \\
= & 1 - Pr_{\Pi_A(q)}
\end{aligned}$$

This proves the lemma. $\square$

Lemma 4.14. *Consider a query φ and let A be an attribute of one of the relations in φ . If $\Pi_A(\varphi)$ has a safe plan, φ also has a safe plan.*

Proof. Let P be a safe plan for query $\Pi_A(\varphi)$ and consider the project operator op in P that projects out A . Create a new plan P' with op removed from P . We will now show that this plan is safe for φ . Since P is a safe plan, every project operator satisfies equation (2). For the operators in P' that are not ancestors of op , Equation (2) remains same. So they are still safe. For the operators in P' that are ancestors of op , A gets added to the left of equation (2). So the functional dependencies still hold and the operators are still safe. Thus, we have shown that a safe plan exists for φ . $\square$

Lemma 4.15. *Consider a plan $P_1 = \Pi_A(q_1 \bowtie q_2)$, where q_1 and q_2 are queries and $Heads(q_1) \subseteq A$. Also, let A contains all the attributes used in the join between q_1 and q_2 . Let $P_2 = q_1 \bowtie (\Pi_{A \cap Heads(q_2)}(q_2))$. Then P_1 is safe $\implies P_2$ is safe.*

Proof. It is easy to observe that P_2 is a valid plan because A contains all the attributes that are used in the final join.

Since $\Pi_A(q_1 \bowtie q_2)$ is safe, the following can be inferred from $\Gamma^p(q_1 \bowtie q_2)$ for each $R^p \in PRel(q_1 \bowtie q_2)$:

$$A, R^p.E \rightarrow Heads(q_1 \bowtie q_2)$$

The above can be rewritten as

$$\begin{aligned}
& Heads(q_1), A \cap Heads(q_2), R^p.E \rightarrow \\
& Heads(q_1) \cup Heads(q_2)
\end{aligned} \tag{3}$$

We know that plans q_1 and q_2 are safe. Thus, to show that P_2 is safe, we only need to show that the project operator $\Pi_{A \cap \text{Heads}(q_2)}(q_2)$ is safe. So, for each $R^p \in \text{PRels}(q_2)$, $\Gamma^p(q_2)$ must imply the following:

$$A \cap \text{Heads}(q_2), R^p.E \rightarrow \text{Heads}(q_2) \quad (4)$$

Let Γ_{join} be the set of FDs introduced by the final join. Then, $\Gamma^p(q_1 \bowtie q_2) = \Gamma^p(q_1) \cup \Gamma^p(q_2) \cup \Gamma_{\text{join}}$. Since A contains all the attributes that occur in Γ_{join} , Γ_{join} is not required for Equation (3). Thus, $\Gamma^p(q_1) \cup \Gamma^p(q_2)$ imply Equation (2) and hence the following:

$$\text{Heads}(q_1), A \cap \text{Heads}(q_2), R^p.E \rightarrow \text{Heads}(q_2)$$

But now, since Γ_{join} is not present, $\text{Heads}(q_1)$ and $\Gamma^p(q_1)$ have no contribution in the above equation. Thus, $\Gamma^p(q_2)$ alone implies Equation (4).

This shows that P_2 is a safe plan. □

Theorem 4.16. *Let φ be a query that has a separation (S, T) . Then, if φ has a safe plan, there must be another safe plan of the form $P_S \bowtie P_T$ such that $\text{PRels}(P_S) = S$ and $\text{PRels}(P_T) = T$ (where the join may contain a set of select conditions).*

Proof. We will prove this by induction on the size of the query. The base step corresponding to queries over a single relation holds trivially. Let P be a safe plan for φ . There are three cases:

1. The top operator on P is a join. Thus, P can be written as $P_1 \bowtie P_2$. Let φ_1 and φ_2 be the queries corresponding to the plans P_1 and P_2 . It is easy to see that $(\text{PRels}(P_1) \cap S, \text{PRels}(P_1) \cap T)$ form a separation for φ_1 . By induction, there is a safe plan $P_{S1} \bowtie P_{T1}$ for φ_1 where $S1 = \text{PRels}(P_1) \cap S$ and $T1 = \text{PRels}(P_1) \cap T$. Similarly, there is a safe plan $P_{S2} \bowtie P_{T2}$ for φ_2 where $S2 = \text{PRels}(P_2) \cap S$ and $T2 = \text{PRels}(P_2) \cap T$. Thus, the following is a safe plan for φ :

$$(P_{S1} \bowtie P_{S2}) \bowtie (P_{T1} \bowtie P_{T2})$$

The safety of the about plan follows from the safety of the individual subplans and the safety of the join operators.

2. The top operator on P is a select. Thus, P can be written as $\sigma_c(P')$. Let φ' be the query corresponding to the plan P' . (S, T) is still a separation for φ' . By induction, there is a safe plan $P'_S \bowtie P'_T$ for φ' . Adding the select condition to the outer join, we get a safe plan as required.
3. The top operator on P is a project. Thus, P can be written as $\Pi_A(P')$. Let φ' be the query corresponding to P' . Note that (S, T) is still a separation for φ' . By induction hypothesis, there is a safe plan for φ' of the form $P'_S \bowtie P'_T$. Also, A can be written as the disjoint union of A_S and A_T ,

where A_S are attributes of relations from S and A_T are attributes of T . Thus, the following plan is safe for φ :

$$\Pi_{(A_S, A_T)}(P'_S, P'_T)$$

Using Lemma 4.13, we get the following equivalent safe plan:

$$\Pi_{A_S, A_T}(\Pi_{Head(P'_S) \cup A_T}(P'_S \bowtie P'_T))$$

We observe that Lemma 4.15 is applicable to the plan $(\Pi_{Head(P'_S) \cup A_S}(P'_S \bowtie P'_T))$ since (S, T) is a separation. Thus, the following equivalent plan is also safe:

$$\Pi_{A_S}(P'_S \bowtie (\Pi_{A_T}(P'_T)))$$

Another application of Lemma 4.15 yields the following plan

$$(\Pi_{A_S}(P'_S)) \bowtie (\Pi_{A_T}(P'_T))$$

Setting $P_S = \Pi_{A_S}(P'_S)$ and $P_T = \Pi_{A_T}(P'_T)$, we get the required result.

The theorem then follows from the principle of induction. $\square$

Theorem 4.17. *Algorithm SAFE-PLAN is complete.*

Proof. Let q be the given query. Suppose it has an attribute A belonging to $Attr(q) - Head(q)$ such that the operator $\Pi_{Head(q)}(q_A)$ is a safe operator, where q_A is as defined in the algorithm. Then, by Lemma 4.14, q has a safe plan if and only if q_A has a safe plan. So the algorithm recursively solves for q_A . On the other hand, suppose there is no such attribute A . This implies that in any safe plan, the final operator cannot be a project operator as it is not safe. So a safe plan must consist of a join of two subplans. These two subplans must form a separation because after they are joined, there is no projection. Hence, there must exist a separation for q . If not, the algorithm correctly returns false. Otherwise, by Theorem 4.16, it is sufficient to solve each of the subproblems corresponding to the separation separately. This completes the proof of correctness of the algorithm. $\square$

5 Theoretical Analysis

We show here a fundamental result on the complexity of query evaluation on probabilistic databases. It forms a sharp separation of conjunctive queries into queries with low and high data complexity, and shows that our optimization algorithm is complete.

The data complexity of a query q is the complexity of evaluating $q^{rank}(D^p)$ as a function of the size of D^p . If q has a safe plan p , then its data complexity is in PTIME, because all extensional operators are in PTIME. We start by showing that, for certain queries, the data complexity is $\#P$ -complete. $\#P$ is the complexity class of some hard counting problems. Given a boolean formula φ , counting the number of satisfying assignments, denote it $\#\varphi$, is $\#P$ -complete [24].

(Checking satisfiability, $\#\varphi > 0$, is NP-complete.) The data complexity of any conjunctive query is $\#P$, since $q^{rank}(D^p) = Pr(q^i(D^p))$. The following is a variant of a result on query reliability by Gradel et al. [12]. The proof is novel and is of independent interest in our setting.

Theorem 5.1. *Consider the following conjunctive query on three probabilistic tables:*

$$q() := L^p(x), J(x, y), R^p(y)$$

Here L^p, R^p are extensional probabilistic tables and J is deterministic⁷. The data complexity for q is $\#P$ -hard.

Proof. (Sketch) Provan and Ball [20] showed that computing $\#\varphi$ is $\#P$ -complete even for *bipartite monotone 2-DNF* boolean formulas φ , i.e. when the propositional variables can be partitioned into $X = \{x_1, \dots, x_m\}$ and $Y = \{y_1, \dots, y_n\}$ s.t. $\varphi = C_1 \vee \dots \vee C_k$ where each clause C_i has the form $x_j \wedge y_k$, $x_j \in X, y_k \in Y$. (The satisfiability problem, $\#\varphi > 0$, is trivially true.). Given φ , construct the instance D^p where L^p is X , R^p is Y and J is the set of pairs (x_j, y_k) that occur in some clause C_i . Assign independent probability events to tuples in L^p, R^p , with probabilities $1/2$. Then $q^{rank}(D^p)$ returns a single tuple, with probability $\#\varphi/2^{m+n}$. Thus, computing $q^{rank}(D^p)$ is at least as hard as computing $\#\varphi$. $\square$

We state now the main theoretical result in this paper. We consider it to be a fundamental property of query evaluation on probabilistic databases.

Theorem 5.2 (Fundamental Theorem of Queries on Probabilistic DBs). *Consider a schema $\bar{R}^p, \Gamma^p$ where all relations are probabilistic and extensional. $Attrs(R) \rightarrow R^p.E$, $R^p.E \rightarrow Attrs(R)$, for every R^p . Let q be a conjunctive query s.t. each relation occurs at most once. Assuming $\#P \neq PTIME$ the following statements are equivalent:*

1. *The query q contains three subgoals R , L and J , and two variables x and y such that:*
 - *J determines x and y , i.e., the variables in J together with the FDs induced by q , determine x and y .*
 - *L determines x but not y .*
 - *R determines y but not x .*
2. *The data complexity of q is $\#P$ -complete.*
3. *The SAFE-PLAN optimization algorithm fails to return a plan.*

Theorem 5.2 provides a sharp separation of feasible and infeasible queries on probabilistic databases. To illustrate the proof of theorem, we prove it for the following special case.

⁷Allowing J to be deterministic strengthens the result. The theorem remains true if J is probabilistic.

Theorem 5.3. Consider a schema $\bar{R}^p, \Gamma^p$ where all relations are probabilistic and Γ^p has only the trivial FDs⁸ $\text{Attrs}(R) \rightarrow R^p.E$, $R^p.E \rightarrow \text{Attrs}(R)$, for every R^p . Let q be a conjunctive query s.t. each relation occurs at most once. Assuming $\#P \neq PTIME$ the following statements are equivalent:

1. The query q contains three subgoals of the form⁹:

$$L(x, \dots), J(x, y, \dots), R(y, \dots)$$

where $x, y \notin \text{Head}(q)$.

2. The data complexity of q is $\#P$ -complete.
3. The SAFE-PLAN optimization algorithm fails to return a plan.

Before we give the proof, we need the following Lemma.

Lemma 5.4. Consider a schema $\bar{R}^p, \Gamma^p$ as described in Theorem 5.2 and q be any conjunctive query. Then, $Pi_{\text{Heads}(q)-x}(q)$ is safe if every relation in $\text{Rels}(q)$ contain x .

Proof. (Lemma 5.4) Let us assume that every relation in $\text{Rels}(q)$ contain R . We will prove that the project operator is safe. We need to show that for each R , the following holds

$$\text{Heads}(q) - x, R.E \rightarrow \text{Heads}(q)$$

From schema, we have the following FD:

$$R.E \rightarrow \text{Attrs}(R)$$

Also, since x is in all relations, we have

$$\text{Attrs}(R) \rightarrow x$$

. Thus,

$$\text{Heads}(q) - x, R.E \rightarrow \text{Heads}(q) - x, x \rightarrow \text{Heads}(q)$$

□

Proof. (Theorem 5.2)

(1) $\Rightarrow$ (2) is a simple extension of Th. 5.1. Let φ be the query as in the proof of Th. 5.1 with X and Y as the set of variables. We construct a database as follows: Consider the canonical database for the query q . For each relation that contains x but not y , replicate the canonical tuple in that relation $|X|$ times with x replaced by the various values in X . For relations containing y but not x , replicate the canonical tuple with y taking values from Y . For relations

⁸Hence, the probabilistic instances are extensional.

⁹Here we are using an alternate representation of conjunctive queries where all occurrences of variables that are equated are replaced by the same symbol

containing both x and y , replicate the canonical tuple once for each clause $(x_i \wedge y_k)$ with x and y replaced with x_i and y_k . Also, set the probabilities of all tuples to 1 except for relations L and R where the tuple probabilities are 0.5. It is easy to see that the answer contains a single tuple which is the canonical answer and its probability is $\# \varphi / 2^{m+n}$. This reduction shows that evaluating q is $\# - P$ complete.

(2) $\Rightarrow$ (3) is obvious, since any safe plan has data complexity in PTIME.

Finally, we prove that (3) $\Rightarrow$ (1). So let us suppose that SAFE-PLAN algorithm fails to return a plan. The only way this can happen is when the algorithm finds a sub-query that has no separation. So let q' be the sub-query of q that has no separation. Let $A_P = \text{Head}(q') - \text{Attrs}(q')$ be the set of attributes that are projected out. So, none of the attributes in A_P is safe to project out. Otherwise, the algorithm will not stop at this query. Since the query has no separation, any pair of relations in $\text{Rels}(q')$ share an attribute from A_P . Also, by Lemma 5.4, no attribute in A_P is contained in all the relations in $\text{Rels}(q')$. Let x be the attribute that is contained in largest number of relations in $\text{Rels}(q')$ and let S_x be the set of relations containing x . Thus, $S_x \subsetneq \text{Rels}(q')$. Some relation in S_x , say J , must share an attribute with some relation outside S_x , say R . Let the common attribute be y (it cannot be x because $R \notin S_x$). Also, there must be a set L in S_x that does not contain y (otherwise, number of sets containing y will be strictly greater than the number of sets containing x). Thus, we have found sets L , J and R such that L contains x but not y , R contains y but not x , and J contains both. This completes the proof. $\square$

Theorem 5.2 states that for the schemas that satisfy certain properties, unsafe queries are $\#P$ -complete. The following theorem generalizes Theorem 5.2 (by removing the requirement that all relations be probabilistic).

Theorem 5.5. *Consider a schema $\bar{R}^P, \Gamma^P$ where all relations are extensional. Let q be a conjunctive query s.t. each relation occurs at most once. Then, if q is not safe, the data complexity of q is $\#P$ -complete.*

Proof. It is sufficient to prove the theorem only for those queries whose head is empty. To see this, if q is any general query that is not safe, then $q' = \Pi_\phi(q)$ is also not safe. Also, the problem of evaluating q can be reduced to the problem of evaluating q' as follows. Consider each tuple t in the output of q . We need to find the probability of t . Construct a new database by removing all the tuples whose values do not match with t . Evaluate q' on this new database. It is easy to see that the result is a single tuple t and its probability is same as that in q . Since there are only polynomial number of tuples in the output of q , this is a polynomial time reduction.

So we assume that $\text{Heads}(q)$ is empty. Since it is not safe, consider a sub-query q_s that has no separation. Let U be the set of unsafe projects. For each relation r , let D_r be the set of attributes x such that $(\text{Attrs}(r) \rightarrow x)$ can be inferred from $\Gamma^P(q)$. If x is in D_r , we say that r determines x . For each attribute in U , there is at least one probabilistic relation that does not determine that

attribute (otherwise, it would be a safe project). Let a be the attribute that is determined by the largest number of relations and let R_a be the set of those relations. There is at least one probabilistic relation not contained in R_a . Let $r_1, r_2 \dots r_{k+1}$ be the smallest sequence of relations satisfying the following: (1) $r_1 \in R_a$, (2) $r_i \notin R_a (i > 1)$, (3) $D_{r_i} \cap D_{r_{i+1}} \neq \phi$ and (4) r_{k+1} is probabilistic. There is at least one such sequence because there is at least one probabilistic relation outside R_a and it must be connected to some relation in R_a (as the query has no separation). Since b be any attribute in $D_k \cup D_{k+1}$. Thus, there is at least one relation $r_0 \in R_a$ that does not determine b (otherwise $|R_b|$ would be strictly greater than $|R_a|$).

We will now show a reduction from the problem of counting the number of assignments of *bipartite monotone 2-DNF*. Let $X = \{x_1, \dots, x_m\}$, $Y = \{y_1, \dots, y_n\}$ and $\varphi = C_1 \vee \dots \vee C_k$ where each clause C_i has the form $x_{f(i)} \wedge y_{f(i)}$. We will construct a database as follows. Consider a single table whose attributes are all the variables that occur in the query. Let $A = D_{r_0}$, $B = D_{r_{k+1}}$, $I = A \cup B$ and Z be the set of remaining variables. Populate the table with k tuples, one corresponding to each clause. In row i , assign a value of $x_{f(i)}$ to all variables in $A - B$, a value of $y_{f(i)}$ to variables in $B - A$, a constant value 0 to variables in $A \cup B$ and a value of C_i to variables in Z . From this table, construct individual tables by taking the projections on corresponding columns. For tables r_0 and $r_{r_{k+1}}$, let all events have a probability of 0.5 (which is valid because both of them are probabilistic) and let all other tables have events with probability 1. To see that all the functional dependencies are satisfied, consider any FD, $X \rightarrow y$, where X is a set of attributes and y is an attribute. For this to be violated, either $X \subseteq A$ and $y \in B - A$ or $X \subseteq B$ and $y \in A - B$. None of these can happen by definitions of A and B . Thus, resulting database does indeed satisfy the schema.

We will now show that the number of satisfying assignments to the DNF is the probability of the binary query q on this database times 2^{m+n} . To show this, it is enough to prove that when all the tables are joined, we get back the universal table, i.e. an x_i in r_0 is joined to a y_j in r_{k+1} only if (x_i, y_j) is a clause. To prove this, consider the relations $r_1, r_2 \dots r_k$. We have $D_{r_i} \cup B = \phi (1 < i < k)$ (otherwise, we would have got a shorter sequence $r_1, r_2 \dots r_i, r_{k+1}$ satisfying all the above properties). Similarly, $D_{r_i} \cup A = \phi$. Thus, for $1 < i < k$, $D_{r_i} \subseteq Z$. Thus, the sub-query q_J consisting of $r_1, r_2, \dots r_k$ generates exactly k tuples, one for each clause. Also, we have $D_{r_0} \cup D_{r_1} \neq \phi$ and $D_{r_k} \cup D_{r_{k+1}} \neq \phi$. Thus, q_J acts as a join between r_0 and r_{k+1} and only joins those variables that correspond to a clause.

This completes the reduction and proves that the query is $\#P$ -complete. $\square$

6 Unsafe Plans

When a query's data complexity is $\#P$ -complete, then SAFE-PLAN fails to return a plan. Since this can indeed happen in practice, we address it and propose two solutions.

6.1 Least Unsafe Plans

Here we attempt to pick a plan that is less unsafe than others, i.e. minimizes the error in computing the probabilities. Recall from Eq.(2) that $\Pi_{A_1, \dots, A_k}^e$ is safe in $\Pi_{A_1, \dots, A_k}^e(q)$ iff $A_1, \dots, A_k, R^p.E \rightarrow Head(q)$ for every R^p . Let $\bar{B} = \{A_1, \dots, A_k, R^p.E\} \cap Attr(R^p)$ (hence $R^p.E \in \bar{B}$) and $\bar{C} = Head(q) \cap Attr(R^p)$. Define R_{fanout}^p to be the expected number of distinct values of $\bar{C}$ for a fixed value of the attributes $\bar{B}$. In a relational database system, it is possible to estimate this value using statistics on the table R^p . Define the degree of unsafety of $\Pi_{A_1, \dots, A_k}^e$ to be $\max_{R^p \in PREL(Q)} (R_{fanout}^p - 1)$. Thus, a safe project has degree of unsafety 0. Also, the higher the degree of unsafety, the higher is the expected error that would result from using the extensional semantics for that project operator.

We modify Algorithm 1 to cope with unsafe queries. Recall that the algorithm tries to split a query q into two subqueries q_1, q_2 s.t. all their join attributes are in $Head(q)$. Now we relax this: we allow joins between q_1 and q_2 on attributes not in $Head(q)$, then project out these attributes. These projections will be unsafe, hence we want to minimize their degree of unsafety. To do that, we pick q_1, q_2 to be a minimum cut of the graph, where each edge representing a join condition is labeled with the degree of unsafety of the corresponding project operation¹⁰. The problem of finding minimum cut is polynomial time solvable as a series of network flow problems or using the algorithm of Stoer and Wagner [21].

6.2 Monte-Carlo Approximations

As an alternative, we present now an algorithm based on a Monte-Carlo simulation, which can guarantee an arbitrarily low error.

Given a conjunctive query q over probabilistic relations $R_1^p, R_2^p \dots R_k^p$, let q' be its body, i.e. $Head(q') = Attr(q') = Attr(q)$ and $q = \Pi_{Head(q)}(q')$. Modify q' to also return all event attributes $\bar{E} = R_1^p.E, \dots, R_k^p.E$. Evaluate q' over the probabilistic database, and group of tuples in the answer based on the values of their attributes $Head(q)$. Consider one such group, and assume it has n tuples $t_1, \dots, t_n$. The group defines the following complex event expression: $\bigvee_{i=1}^n C_i$, where each C_i has the form $e_1 \wedge \dots \wedge e_k$. We need to compute its probability, since this will be the probability of one tuple in $q^{rank}(D^p)$. For that we use the Monte Carlo algorithm described by Karp [15]: when run for $N \geq \frac{4n}{\epsilon^2} \ln \frac{2}{\delta}$ iterations, the algorithm guarantees that the probability of the error being greater than ϵ is less than δ .

7 Extensions

Additional operators So far, we have limited our discussion to conjunctive

¹⁰The estimator of R_{fanout}^p should make sure that the estimated value is 0 only when the FD holds, otherwise the algorithm may favor ‘expected’ safe plans over truly safe plans.

queries, or, equivalently to the algebra consisting of $\sigma, \Pi, \times$. We show now how to extend these techniques to $\cup, -, \gamma$ (union, difference, groupby-aggregate). A large fragment of SQL, including queries with nested sub-queries, aggregates, group-by and existential/universal quantifiers can be expressed in this logical algebra [23]. (We omit δ (duplicate elimination) since we only consider queries with set semantics, i.e. δ is implicit after every projection and union.) We define the extensional semantics for these operators, using the functional notation.

$$\begin{aligned}
Pr_{p \cup^e p'}(t) &= 1 - (1 - Pr_p(t)) \times (1 - Pr_{p'}(t)) \\
Pr_{p -^e p'}(t) &= Pr_p(t) \times (1 - Pr_{p'}(t)) \\
Pr_{\gamma_{\bar{A}, \min(B)}^e}(t) &= Pr_p(t) \times \prod_{\substack{s : s.\bar{A} = t.\bar{A} \\ \wedge s.B < t.B}} (1 - Pr_p(s)) \\
Pr_{\gamma_{\bar{A}, \max(B)}^e}(t) &= Pr_p(t) \times \prod_{\substack{s : s.\bar{A} = t.\bar{A} \\ \wedge s.B > t.B}} (1 - Pr_p(s))
\end{aligned}$$

For example, to compute the groupby-min operator $\gamma_{A, \min(B)}(R^p)$ one considers each tuple (a, b) in R^p : the probability that (a, b) is in the output relation is $p(1 - p_1) \dots (1 - p_n)$ where p is the probability of the tuple (a, b) , while $p_1, \dots, p_n$ are the probabilities of all other tuples (a, b') s.t. $b' < b$. In the case of SUM, the aggregated attribute may take values that are not in the input table. To compute the probabilities correctly one needs to iterate over exponentially many possible sums. Instead, we simply compute the expected value of the sum (details omitted). This is meaningful to the user if SUM appears in the SELECT clause, less so if it occurs in a HAVING clause. We treat COUNT similarly.

We now give sufficient conditions for these operators to be safe.

Theorem 7.1. *Let q, q' be conjunctive queries.*

1. $\cup^e$ is safe in $q \cup^e q'$ if $PRels(q) \cap PRels(q') = \phi$.
2. $-^e$ is safe in $q \cap^e q'$ if $PRels(q) \cap PRels(q') = \phi$.
3. $\gamma_{\bar{A}, agg(B)}^e$ is safe in $\gamma_{\bar{A}, agg(B)}(q)$ if $\Pi_{\bar{A}}(q)$ is safe, where agg is min or max.

Self-joins Self-joins on probabilistic relations may be a cause of $\# - P$ -complete data complexity [12]. However, a query $q^\approx$ with uncertain predicate rarely results in self-join. Even if the same table R occurs twice in $q^\approx$, the different uncertain predicates on the two occurrences generate distinct events, hence the system makes two probabilistic “copies”: R_1^p, R_2^p . Of course, the Monte-Carlo algorithm works fine even in the presense of self-joins.

Extending the optimization algorithm SAFE-PLAN is extended to handle each block of conjunctive queries separately. As an example, the query in

Section 1, asking for an actor whose name is like ‘Kevin’ and whose first ‘successful’ movie appeared in 1995, has a safe plan as shown below:

$$\Pi_{name}(A \bowtie_{actorid} (\sigma_{year=1995}(\gamma_{actorid, \min(year)}(\Pi_{actorid, year} C))))$$

8 Atomic Predicates

Our main motivation is executing a query with uncertain predicates $q^\approx$ on a deterministic database D . As we saw, our approach is to apply the uncertain predicates first, and generate a probabilistic database D^p , then evaluate q (without the uncertain predicates). We discuss here briefly some choices for the uncertain predicates proposed in the literature. All proposals depend on a notion of closeness between two data values. This is domain dependent and can be classified into three categories:

Syntactic closeness This applies to domains with proper nouns, like people’s names. Edit distances, q-grams and phonetic similarity can be employed. The excellent surveys on string matching techniques by Zobel and Dart [25] and Navarro [18] describe more than 40 techniques and compare them experimentally. Navarro also has a discussion on the probability of string matching. In our system, we used the 3-gram distance between words, which is the number of triplets of consecutive words common to both words. We found it to work faster and more accurately than edit distances.

Semantic closeness This applies to domains that have a semantic meaning, like film categories. A user query for the category ‘musical’ should match films of category ‘opera’. Semantic distance can be calculated by using TF/IDF or with ontologies like Wordnet [2]. We do not have semantic distances in our system currently.

Numeric closeness This applies to domains like *price* and *age*. A distance can be just the difference of the values.

Once distances are defined between attributes, using any of the above methods, they need to be meaningfully converted into probabilities. We fitted a Gaussian curve on the distances as follows: the curve was centered around the distance 0 where it took value 1. The variance of the Gaussian curve is an indication of the importance of match on that attribute. Its correct value depends on the domain and user preferences. In our experiments, we used fixed, query independent values, for the variances.

Finally, one issue is when to generate new probability events. For example consider the uncertain predicate $\mathbf{Product.category} \approx \dots$ and assume there are two products with the same category. Should they result in two independent probabilistic events with the same probabilities, or in the same probabilistic events? Both choices are possible in our system. In the first case the functional dependency is $\mathbf{Product}^p.\mathbf{key} \rightarrow \mathbf{Product}^p.E$ while in the second the FD

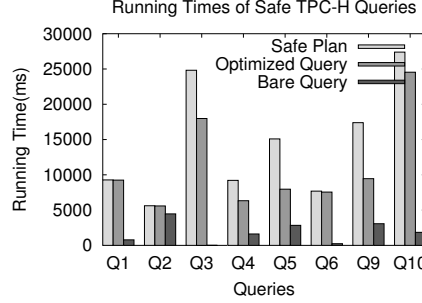


Figure 7: TPC-H Query Running Times

is $\text{Product}^p.\text{category} \rightarrow \text{Product}^p.E$. In the latter case, $\Pi_{category}$ becomes unsafe. This can be taken care of by normalizing the resulting database to 3NF, i.e. creating a separate category table that contains the events for categories.

9 Experiments

We performed some preliminary evaluation of our probabilistic query evaluation framework, addressing four questions. How often does the SAFE-PLAN optimization algorithm fail to find a plan? What is the performance of safe plans, when they exist? Are naive approaches to query evaluation perhaps almost as good as a safe plan? And how effectively can we handle queries that do not have safe plans?

We did not modify the relational engine, but instead implemented a middleware. SQL queries with approximate predicates were reformulated into “extensional” SQL queries, using the techniques described in this paper, and calls to a TSQL function computing 3-gram distances. These queries were then executed by the relational engine and returned both tuples and probabilities. We used Microsoft SQL Server.

We used the TPC-H benchmark, with a database of 0.1GB. We modified all queries by replacing all the predicates in the WHERE clause with uncertain matches. The constants in the queries were either misspelled or made vague. For instance, a condition like `part.container = 'PROMO PLATED GREEN'` was replaced with `part.container ≈ 'GREEN PLATE'`. When executed exactly, all modified queries returned empty answers.

1. Frequency of unsafe queries In our first experiment, we wanted to see how many queries do not have safe plans. Out of the 10 TPC-H queries, 8 turned out to have safe plans. Q_7 and Q_8 were the only query that were unsafe. These also become safe if not all of their predicates are uncertain.

2. Performance Next, we measured the running times for the eight queries that have safe plans, shown in Figure 7. All times are wall-clock. The first column is the running time of the safe plan. The second column represents an

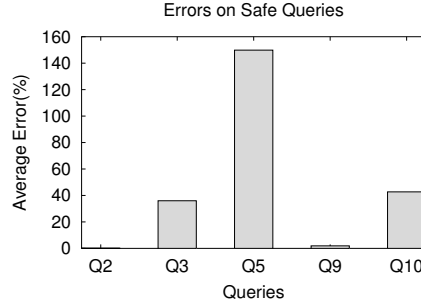


Figure 8: Errors on Safe TPC Queries

optimization where at each intermediate stage, tuples with zero probability are discarded. This optimization does not affect the final answer and as we can see from the graph, it brings about considerable savings for some queries. This also suggests the use of other optimizations like an early removal of tuples with low probabilities if the user is only interested in tuples with high probability. The third column in the graph shows the time for running safe queries without taking into account the computation time for the uncertain predicate, which, in our case, is the 3-gram distance. The graphs show that most of the time is spent in computing the uncertain predicates. (For Q_3 the the running time was almost negligible.) This graph suggests that important improvements would be achieved if the predicates were implemented in the engine.

3. Naive Approaches In the next experiment we calculated the error produced by a naive extensional plan. We considered the naive plan that leaves all project operators (and the associated duplicate elimination) at the end of the plan, which are typical plans produced by database optimizers. Figure 8 shows the percentage relative error of naive plans. We only considered the 8 queries that have safe plans. The naive plans for Q_1 , Q_4 , Q_6 were already safe, hence had no errors (and SAFE-PLAN indeed returned the same plan): these queries are not shown. Queries Q_3 , Q_5 and Q_{10} had large errors with Q_5 showing an average error of 150% in the tuple probabilities. Queries Q_2 and Q_9 had negligible errors. Thus, while some naive plans were bad, others were reasonable. But, in general, naive plans *can* be arbitrarily bad. However, we argue that the low extra complexity of searching for a safe plan is a price worth paying in order to avoid the (admittedly rare) possibility of arbitrarily large errors.

However, since we are only interested in *ranking* the results, not in the actual probabilities, it is worth asking whether high errors in the probabilities translate into high ranking results. We plotted the recall graphs for queries Q_3 and Q_{10} (for which the naive plan produced only medium errors). We defined recall as the fraction of answers ranked among top N by the naive plan that should actually have been in top N . We plotted this as a function of N . Figures 9

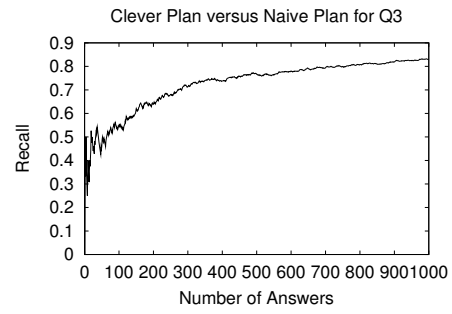


Figure 9: Recall Plot for Q_3

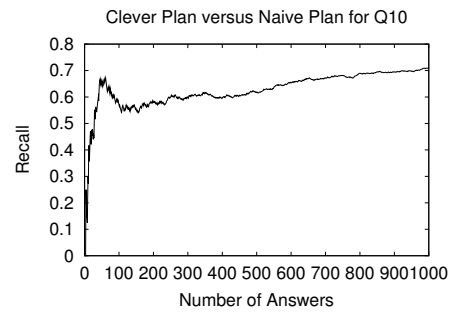


Figure 10: Recall Plot for Q_{10}

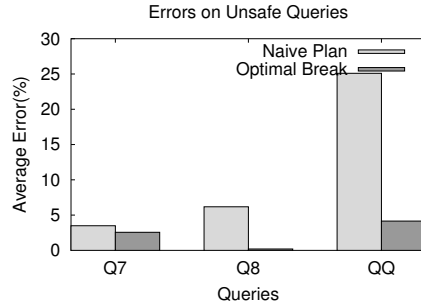


Figure 11: Errors on Unsafe Queries

and 10 show the recall graphs. By definition, the recall approaches to 1 when N approaches the total number of possible tuples in the answer. However, as the graphs show, the recall was bad for small values of N . A user looking for top 50 or 100 answers to Q_3 would miss half of the relevant tuples. For smaller values of N (say, 10) the naive approach misses 80% of the relevant tuples.

4. Unsafe Queries Finally, we tested our approach to handle queries with no safe plans on Q_7 and Q_8 . We ran the Monte Carlo simulation to compute their answer probabilities and used them as baseline. Figure 11 shows the errors in evaluating them with a naive plan and the least unsafe plan (using min-cut, Sec. 6). The graphs show that the plan chosen by the optimizer was better, or significantly better than a naive one. Still, from two data points it is hard to judge the improvement over a naive plan. To see a third data point we wrote a new unsafe query, QQ , where the relation `lineitem` is joined with `orders` and `suppliers`. Here the fanout is larger, and the difference between the naive plan and the optimal break is more pronounced.

10 Related Work

The possible worlds semantics, originally put forward by Kripke for modal logics, is commonly used for representing knowledge with uncertainties. Halpern, Baccus et al [9, 4] have showed the use of possible worlds semantics to assign degrees of beliefs to statements based of the probabilities in the knowledge base.

Though there has been extensive work on probabilistics in AI, relatively little work has been done on probabilistic databases. There are probabilistic frameworks [7, 6, 11, 16] proposed for databases, but each makes simplifying assumptions for getting around the problem of high query evaluation complexity that lessens their applicability.

Fuhr and Rolke [11] define probabilistic NF2 relations and introduce the intensional semantics for query evaluation. As we saw, this is correct, but impractical.

Many of these works specialize on logic programming in deductive databases. Ng and Subrahmaniam [19] extend deductive databases with probabilistics and give fixed point semantics to logic programs annotated with probabilities, but they use absolute ignorance to combine event probabilities.

Non-probabilistic approaches to imprecise queries have also been considered. Keyword searches in databases are discussed in [14, 8, 13]. Fagin [10] gives an algorithm to rank objects based on its scores from multiple sources: this applies only to a single table. The VAGUE system [17] supports queries with vague predicates, but the query semantics are ad hoc, and apply only to a limited SQL fragments. Surajit et al. [3] consider ranking query results automatically: this also applies to a single table. Theobald and Weikum [22] describe a query language for XML that supports approximate matches with relevance ranking based on ontologies and semantic similarity.

11 Conclusions

In this paper, we introduce a query semantics on probabilistic databases based on *possible worlds*. Under this semantics, every query that is well defined over a deterministic databases has a meaning on a probabilistic database. We describe how to evaluate queries efficiently under this new semantics. Our theoretical results capture fundamental properties of queries on probabilistic databases, and lead to efficient evaluation techniques. We showed how this approach can be used to evaluate arbitrarily complex SQL queries with uncertain predicates.

References

- [1] Movie database: <http://kdd.ics.uci.edu/database-s/movies/movies.html>.
- [2] Wordnet 2.0: A lexical database for the english language: <http://www.cogsci.princeton.edu/wn/>, Jul. 2003.
- [3] S. Agrawal, S. Chaudhuri, G. Das, and A. Gionis. Automated ranking of database query results. In *Proceedings of the First Biennial Conf. on Innovative Data Systems Research*, 2003.
- [4] Fahiem Bacchus, Adam J. Grove, Joseph Y. Halpern, and Daphne Koller. From statistical knowledge bases to degrees of belief. *Artificial Intelligence*, 87(1-2):75–143, 1996.
- [5] R. Baeza-Yates and B. Ribeiro-Neto. *Modern Information Retrieval*. Addison-Wesley, 1999.
- [6] Daniel Barbará, Hector Garcia-Molina, and Daryl Porter. The management of probabilistic data. *IEEE Trans. Knowl. Data Eng.*, 4(5):487–502, 1992.

- [7] Roger Cavallo and Michael Pittarelli. The theory of probabilistic databases. In *VLDB'87, Proceedings of 13th Int. Conf. on Very Large Data Bases, September 1-4, 1987, Brighton, England*, pages 71–81, 1987.
- [8] S. Chaudhuri, G. Das, and V. Narasayya. Dbexplorer: A system for keyword search over relational databases. In *Proceedings of the 18th Int. Conf. on Data Engineering, San Jose, USA*, 2002.
- [9] Ronald Fagin and Joseph Y. Halpern. Reasoning about knowledge and probability. In *Proceedings of the Second Conf. on Theoretical Aspects of Reasoning about Knowledge*, pages 277–293, San Francisco, 1988.
- [10] Ronald Fagin, Amnon Lotem, and Moni Naor. Optimal aggregation algorithms for middleware. In *Proceedings of the twentieth ACM SIGMOD-SIGACT-SIGART symposium on Principles of database systems*, pages 102–113, 2001.
- [11] Norbert Fuhr and Thomas Rolleke. A probabilistic relational algebra for the integration of information retrieval and database systems. *ACM Trans. Inf. Syst.*, 15(1):32–66, 1997.
- [12] Erich Gradel, Yuri Gurevich, and Colin Hirsch. The complexity of query reliability. In *Symposium on Principles of Database Systems*, pages 227–234, 1998.
- [13] Lin Guo, Feng Shao, Chavdar Botev, and Jayavel Shanmugasundaram. Xrank: Ranked keyword search over xml documents. In *Proceedings of the 2003 ACM SIGMOD Int. Conf. on Management of Data, San Diego, California, USA, June 9-12, 2003*, pages 16–27, 2003.
- [14] V. Hristidis and Y. Papakonstantinou. Discover: Keyword search in relational databases. In *Proc. 28th Int. Conf. Very Large Data Bases, VLDB*, 2002.
- [15] Richard Karp and Michael Luby. Monte-carlo algorithms for enumeration and reliability problems. In *Proceedings of the annual ACM symposium on Theory of computing*, 1983.
- [16] Laks V. S. Lakshmanan, Nicola Leone, Robert Ross, and V. S. Subrahmanian. Proview: a flexible probabilistic database system. *ACM Trans. Database Syst.*, 22(3):419–469, 1997.
- [17] Amihai Motro. Vague: a user interface to relational databases that permits vague queries. *ACM Trans. Inf. Syst.*, 6(3):187–214, 1988.
- [18] Gonzalo Navarro. A guided tour to approximate string matching. *ACM Computing Surveys*, 33(1):31–88, 2001.
- [19] Raymond T. Ng and V. S. Subrahmanian. Probabilistic logic programming. *Information and Computation*, 101(2):150–201, 1992.

- [20] J. S. Provan and M. O. Ball. The complexity of counting cuts and of computing the probability that a graph is connected. *SIAM J. Comput.*, 12(4):777–788, 1983.
- [21] M. Stoer and F. Wagner. A simple min cut algorithm. *Algorithms-ESA '94*, pages 141–147, 1994.
- [22] Anja Theobald and Gerhard Weikum. The xxi search engine: ranked retrieval of xml data using indexes and ontologies. In *Proceedings of the 2002 ACM SIGMOD Int. Conf. on Management of data*, pages 615–615, 2002.
- [23] Jeffrey D. Ullman and Jennifer Widom. *First Course in Database Systems*, 2nd ed. Prentice Hall, 1997.
- [24] L. Valiant. The complexity of enumeration and reliability problems. *SIAM J. Comput.*, 8:410–421, 1979.
- [25] J. Zobel and P. W. Dart. Phonetic string matching: Lessons from information retrieval. In *Proceedings of the 19th Int. Conf. on Research and Development in Information Retrieval*, pages 166–172, Zurich, Switzerland, 1996. ACM Press.