

# Efficient Evaluation of HAVING Queries on a Probabilistic Database

Christopher Ré and Dan Suciu

Department of Computer Science and Engineering  
University of Washington  
{chrisre,suciu}@cs.washington.edu

**Abstract.** We study the evaluation of positive conjunctive queries with Boolean aggregate tests (similar to HAVING queries in SQL) on probabilistic databases. Our motivation is to handle aggregate queries over imprecise data resulting from information integration or information extraction. More precisely, we study conjunctive queries with predicate aggregates using MIN, MAX, COUNT, SUM, AVG or COUNT(DISTINCT) on probabilistic databases. Computing the precise output probabilities for positive conjunctive queries (without HAVING) is  $\#P$ -hard, but is in  $P$  for a restricted class of queries called safe queries. Further, for queries without self-joins either a query is safe or its data complexity is  $\#P$ -Hard, which shows that safe queries exactly capture tractable queries without self-joins. In this paper, for each aggregate above, we find a class of queries that exactly capture efficient evaluation for HAVING queries without self-joins. Our algorithms use a novel technique to compute the marginal distributions of elements in a semiring, which may be of independent interest.

## 1 Introduction

We study the complexity of evaluating aggregate queries on probabilistic databases. Our motivation is managing data resulting from integration applications, e.g. object reconciliation [15, 30, 31] and information extraction [5, 14, 17, 19]. Standard approaches require that we eliminate all uncertainty before any querying can begin, which is expensive in both man-hours to perform the integration and in lost revenue due to down time. An alternative approach where data are allowed to be uncertain but we capture uncertainty using probabilities has attracted renewed interest [6, 7, 9, 10, 22, 29]. In such systems, individual tuples are allowed to be incorrect, but aggregations of tuples still provide meaningful information.

In SQL, aggregates come in two forms: *value aggregates* that are returned to the user in the SELECT clause (e.g. the MAX price) and *predicate aggregates* that appear in a HAVING clause (e.g. is the MAX price greater than \$10.00?). In this paper, we focus on positive conjunctive queries with a single predicate aggregate which we call HAVING queries. Prior art [4, 16] has defined a semantic for *value aggregation* that returns the expected value of an aggregate query (e.g. the expected MAX price) and have demonstrated its utility for Decision Support or OLAP style applications. In this paper, we propose a complementary semantic for predicate aggregates inspired by HAVING (e.g. what is the *probability* that the MAX price is bigger than \$10.00?). We illustrate the difference between the approaches with a simple example:

*Example 1.* Consider a probabilistic database with a `Profit` relation that contains the profit forecasted for each item if we continue to sell it:

| Item    | Forecaster | Profit | $P$  |
|---------|------------|--------|------|
| Widget  | Alice      | \$-99K | 0.99 |
|         | Bob        | \$100M | 0.01 |
| Whatsit | Alice      | \$1M   | 1    |

SELECT SUM(PROFIT)  
FROM PROFIT  
WHERE ITEM='Widget'

(a) Expectation Style

SELECT ITEM  
FROM PROFIT  
WHERE ITEM='Widget'  
HAVING SUM(PROFIT) > 0.0

(b) HAVING Style

`Profit(Item;Forecaster;Profit;P)`

`Profit` is an example of BID relation which captures our uncertain about contradictory and incompatible forecasts. Here, we trust Alice's forecast (probability 0.99) more than Bob's (0.01). Prior art [16] considered aggregate queries in the `SELECT` clause such as (a). Their semantic computes the *expected profit*. Using linearity of expectation, the value of query (a) is  $100M * 0.01 + -99K * 0.99 \approx 900K$ . This large value suggests that we should continue selling widgets because we expect to make money. However, if we asked the `HAVING` style query (b), which says: What is the chance that we will make a profit? The answer is only 0.01, which tells us that we should immediately stop selling widgets or risk going out of business.

Prior work [7, 11, 21] has shown that for Boolean conjunctive queries without `HAVING`, computing a query's probability is  $\#\mathcal{P}$ -Complete<sup>1</sup>. Although in general evaluating conjunctive queries on a probabilistic database is hard, there are a class of queries that can be computed efficiently called *safe queries* [7, 21]. In this paper, for each aggregate  $\alpha$ , we find a class of `HAVING` queries called  $\alpha$ -safe that can be evaluated efficiently. Further, we show that there is a dichotomy for queries without self-joins: Either  $Q$  is  $\alpha$ -safe, and has an algorithm in  $\mathcal{P}$ , or  $Q$  is not  $\alpha$ -safe and is  $\#\mathcal{P}$ -Hard.

Our starting observation is that evaluating a query with a value aggregate with aggregate  $\alpha$  on a traditional database can be computed by annotating the database with values from some semiring,  $S_\alpha$ , then computing the annotations returned by a the query by evaluating any algebra plan  $P$  over the semiring, using the rules in [12]. Therefore the output of an aggregate function  $\alpha$  on a probabilistic database is described by a random variable  $s_Q$  with values in  $S_\alpha$ , and a `HAVING` query  $Q$  whose predicate is, say, `COUNT(*) < k`, can be computed over the probabilistic database in two stages: first compute the random variable  $s_Q$ , second apply some *recovery function* that computes the probability  $s_Q < k$ . The cost of this algorithm depends on the space required to represent random variables  $s_Q$ , which is proportional to the set of possible worlds of the probabilistic database and hence is prohibitively high. Our key technical insight is that if the plan  $P$  is a *safe plan* then the random variable  $s_Q$  can be computed using a much more concise representation called a *marginal vector*, because  $P$  can guarantee that the random variables being combined are either independent or disjoint. In addition we need to impose an extra condition on  $P$  to ensure that the recovery function can be computed from the marginal vector representation of  $s_Q$ , and we call this condition plus the safety condition for the plan  $P$ , which depends on  $\alpha$ ,  $\alpha$ -safety.

<sup>1</sup>  $\#\mathcal{P}$  defined by Valiant [28] is the class of functions that contains the problem of counting the number of solutions to  $\mathcal{NP}$ -Hard problems (e.g.  $\#3\text{-SAT}$ ).

**Contributions and Overview** We study conjunctive queries with HAVING predicates where the aggregation function is given by MIN, MAX, COUNT, SUM, AVG or COUNT(DISTINCT) and the aggregate test is one of  $=, \neq, <, \leq, >, \geq$  on common representations of probabilistic databases [3, 22, 29]. In Sec. 2, we formalize HAVING queries, our choice of representation and define efficient evaluation. In Sec. 3, we review the relevant technical material (e.g. semirings and safe plans). In Sec. 4, we give our main results: For each aggregate  $\alpha$ , we find a class of HAVING queries, called  $\alpha$ -safe, such that for any  $Q$  using  $\alpha$ :

- If  $Q$  is  $\alpha$ -safe then  $Q$ 's data complexity is in  $\mathcal{P}$ .
- If  $Q$  has no self-joins and is not  $\alpha$ -safe then,  $Q$  has  $\#\mathcal{P}$ -hard data complexity.
- It can be decided in polynomial time in the size of  $Q$  if  $Q$  is  $\alpha$ -safe.

## 2 Formal Problem Description

We first define the syntax and semantics of HAVING queries on probabilistic databases and then define the problem of evaluating HAVING queries.

**Semantics** We consider the aggregates MIN, MAX, COUNT, SUM, AVG and COUNT(DISTINCT) as functions on multisets with the obvious semantics.

**Definition 1.** A Boolean conjunctive query is a single rule  $q = g_1, \dots, g_m$  where each  $g_i$  is a positive EDB predicate. A Boolean HAVING query is a single rule:

$$Q[\alpha(y) \theta k] \text{ :- } g_1, \dots, g_n$$

where for each  $i$ ,  $g_i$  is a positive EDB predicate,  $\alpha \in \{\text{MIN, MAX, COUNT, SUM, AVG, COUNT(DISTINCT)}\}$ ,  $y$  is a single variable<sup>2</sup>,  $\theta \in \{=, \neq, <, \leq, >, \geq\}$  and  $k$  is a constant. The set of variables in the body of  $Q$  is denoted  $\text{var}(Q)$ . We assume that  $y \in \text{var}(Q)$ . The conjunctive query  $q = g_1, \dots, g_n$ , is called the *skeleton* of  $Q$ , denoted  $\text{sk}(Q) = q$ . We call  $\theta$  the *predicate test*,  $k$ , the *predicate operand* and a pair  $(\alpha, \theta)$  an *aggregate test*.

Fig. 1(a) shows a SQL query with a HAVING predicate, that asks for all movies reviewed by at least two distinct reviewers. A translation of this query into an extension of our syntax is shown in Fig. 1(b). The translated query is not a Boolean HAVING query because it has a head variable ( $d$ ). In this paper, we discuss only Boolean HAVING queries. However, as is standard, to study the complexity of non-boolean queries, we can substitute constants for head variables. For example, if we substitute ‘M. Ritchie’ for  $d$ , then the result is Fig. 1(c) which is a Boolean HAVING query.

**Definition 2.** Given a HAVING query  $Q[\alpha(y) \theta k]$  and a relational instance  $W$ , let

$$\mathcal{Y} = \{ \{ v(y) \mid v \text{ is a valuation for } Q \text{ and } \text{im}(v) \subseteq W \} \}$$

i.e.  $\mathcal{Y}$  is the multiset of all valuations of  $Q$  applied to  $y$ . We say that  $Q$  is satisfied on  $W$  and write  $W \models Q[\alpha(y) \theta k]$  (or simply  $W \models Q$ ) if  $\mathcal{Y} \neq \emptyset$  and  $\alpha(\mathcal{Y}) \theta k$  holds.

<sup>2</sup> For COUNT, we will omit  $y$  and write the more familiar COUNT(\*) instead.

|                                                                                                                                                     |                                                                                                                                                           |                                                                                                                                                      |
|-----------------------------------------------------------------------------------------------------------------------------------------------------|-----------------------------------------------------------------------------------------------------------------------------------------------------------|------------------------------------------------------------------------------------------------------------------------------------------------------|
| SELECT m.Title<br>FROM MovieMatch m, Reviewer r<br>WHERE m.ReviewTitle = r.ReviewTitle<br>GROUP BY m.Title<br>HAVING COUNT(DISTINCT r.reviewer) ≥ 2 | $Q(m)[\text{COUNT}(\text{DISTINCT } y) \geq 2] :-$<br>MovieMatch <sup>p</sup> ( <i>t</i> , <i>m</i> ),<br>Reviewer <sup>p</sup> (−, <i>r</i> , <i>t</i> ) | $Q[\text{COUNT}(\text{DISTINCT } y) \geq 2] :-$<br>MovieMatch <sup>p</sup> (‘Fletch’, <i>m</i> ),<br>Reviewer <sup>p</sup> (−, <i>r</i> , <i>t</i> ) |
| (a) SQL Query                                                                                                                                       | (b) Extended Sytanx (Non Boolean)                                                                                                                         | (c) Syntax of this paper                                                                                                                             |

**Fig. 1.** A translation of the query “Which movies have been reviewed by at least 2 distinct reviewers?” into SQL and the syntax of this paper.

**Probabilistic Databases** In this paper, we will use probabilistic databases described in the block-independent disjoint (BID) representation [22, 24] which generalizes many representations in the literature including  $p$ -?-sets and  $p$ -or-sets [13], ?- and  $x$ -relations [26] and tuple independent databases [7, 18] and is similar to [3].

*Syntax* A **BID schema** has relational schemas of the form  $R(K; A; P)$  with its attributes partitioned into three classes separated by semicolons: the **possible worlds key** ( $K$ ), the **value attributes** ( $A$ ), and a single distinguished probability attribute  $P$  taking values in  $(0, 1]$ . The corresponding **possible worlds schema** has relations of the form  $R(K; A)$ , i.e. the same schema without the attribute  $P$ .

*Semantics* Let  $J$  be an instance of a BID schema. We denote by  $t[KAP]$  a tuple in  $J$ , emphasizing its three kinds of attributes, and call  $t[KA]$ , its projection on the  $KA$  attributes, a *possible tuple*. Define a *possible world*,  $W$ , to be any instance consisting of possible tuples s.t.  $K$  is a key in  $W$ . Note that the key constraints do not hold in  $J$ , but do hold in any possible world. Let  $\mathcal{W}_J$  be the set of all possible worlds. We define the semantics of BID instances only for *valid* instances, which are instances  $J$  s.t. for every tuple  $t \in R^J$  in any BID relation  $R(K; A; P)$  the inequality  $\sum_{s \in R: s[K]=t[K]} s[P] \leq 1$  holds. For a valid instance  $J$  its semantics is a finite probability space  $(\mathcal{W}_J, \mu_J)$ . First note that any possible tuple  $t[KA]$  can be viewed as an event in the probability space  $(\mathcal{W}_J, \mu_J)$ , namely the event that a world contains  $t[KA]$ . Then we define the semantics of  $J$  to be the probability space  $(\mathcal{W}_J, \mu_J)$  s.t. (a) the marginal probability of any possible tuple  $t[KA]$  is  $t[P]$ , (b) any two tuples from the same relation  $t[KA]$ ,  $t'[KA]$  s.t.  $t[K] = t'[K]$  are *disjoint events* (a.k.a. exclusive events), and (c) for any set of tuples  $\{t_1, \dots, t_n\}$  s.t. all tuples from the same relation have distinct keys, the events defined by these tuples are independent.

*Example 2.* The data in Fig. 2 shows an example of a BID representation that stores data from integrating extracted movie reviews (e.g. from USENET) with a movie database (e.g. IMDB). The MovieMatch table is uncertain because it is the result of an automatic matching procedure. For example, the probability a review title ‘Fletch’ matches a movie titled ‘Fletch’ is very high 0.95, but not 1 because the title is extracted from text and so may contain errors: For example, from ‘The second Fletch movie’, our extractor will likely extract just ‘Fletch’ although this review actually refers to ‘Fletch 2’. The review table is uncertain because it is the result of *information extraction* and so

| Title              | Matched            | P    |       | ReviewID | Reviewer | Title               | P    |            |
|--------------------|--------------------|------|-------|----------|----------|---------------------|------|------------|
| 'Fletch'           | 'Fletch'           | 0.98 | $m_1$ | 231      | 'Ryan'   | 'Fletch'            | 0.7  | $t_{231a}$ |
| 'Fletch'           | 'Fletch 2'         | 0.9  | $m_2$ |          |          | 'Spies Like Us'     | 0.3  | $t_{231b}$ |
| 'Fletch 2'         | 'Fletch'           | 0.4  | $m_3$ | 232      | 'Ryan'   | 'European Vacation' | 0.90 | $t_{232a}$ |
| 'The Golden Child' | 'The Golden Child' | 0.95 | $m_4$ |          |          | 'Fletch 2'          | 0.05 | $t_{232b}$ |
| 'The Golden Child' | 'Golden Child'     | 0.8  | $m_5$ | 235      | 'Ben'    | 'Fletch'            | 0.8  | $t_{235a}$ |
| 'The Golden Child' | 'Wild Child'       | 0.2  | $m_6$ |          |          | 'Wild Child'        | 0.2  | $t_{235b}$ |

MovieMatch(CleanTitle,ReviewTitle;P)      Reviews(Reviewer,ReviewTitle;Rating;P)

**Fig. 2.** Sample Data arising from integrating automatically extracted reviews from a movie database. **MovieMatch** is a probabilistic relation, we are uncertain which review title matches with which movie in our clean database. **Reviews** is uncertain because it is the result of *information extraction* and *sentiment analysis*.

we have extracted the title from free text (e.g. 'Fletch is a great movie, just like Spies Like Us'). Notice that  $t_{232a}[P] + t_{232b}[P] = 0.95 < 1$ , which indicates that there is some probability reviewid 232 is actually not a review at all.

*Remark 1.* Recall that two distinct possible  $t[KA]$  and  $t'[KA]$  are disjoint if  $t[K] \neq t'[K]$  and  $t[A] = t'[A]$ . But what happens if  $A = \emptyset$ , i.e. all attributes are part of the possible worlds key? In that case all possible tuples become independent, and we sometime call a table  $R(K; ; P)$  a *tuple independent table* [7] or a *?-table* [20] or a *p-?-table* [13].

Queries are posed over the possible worlds schema. For clarity, we denote such relations with a superscripted  $p$  (e.g.  $\text{Reviews}^p$ ).

**Definition 3 (Query Semantics).** The marginal probability of a *HAVING* query  $Q$  on *BID* database  $J$  is denoted  $\mu_J(Q)$  (or simply  $\mu(Q)$ ) and is defined by:

$$\mu_J(Q) = \sum_{W \in \mathcal{W}_J : W \models Q} \mu_J(W)$$

We also make use of  $\mu_J(q)$  where  $q$  is a Boolean conjunctive query with the standard semantics.

*Example 3.* Fig. 1(c) shows a query which asks for all movies that were reviewed by at least 2 different reviewers. The movie 'Fletch' is present when the formula is satisfied  $(m_1 \wedge t_{231a}) \vee (m_2 \wedge t_{232b}) \vee (m_1 \wedge t_{235a})$  and the multiplicity of tuples is exactly the number of disjuncts satisfied. Thus,  $\mu(Q)$  is the probability that at least two of these disjuncts are true, which semantically can be computed by summing over all possible worlds.

**Notions of complexity for HAVING queries** In the database tradition, we would like to measure the data complexity [1], i.e. treat the query as fixed, but allow the data to grow. This assumption makes sense in practice because the running time for query evaluation can be  $O(n^{f(|Q|)})$  where  $|Q|$  is the size of a conjunctive query  $Q$ . This exponential running

time is considered to be tenable in practice, because database queries are generally small. However, in our setting this introduces a problem. By fixing a HAVING query  $q$  we also fix  $k$ , which means that we should accept a running time  $n^k$  as efficient. Clearly this is undesirable: because  $k$  can be large. For example,  $Q()[\text{SUM}(y) > 200] :- R(x, y)$ . For that reason we consider in this paper an alternative definition of the data complexity of HAVING queries, where both the database and  $k$  are part of the input.

**Definition 4.** Fix a skeleton  $q$ , an aggregate  $\alpha$ , and a comparison operator  $\theta$ . The **query evaluation problem** is: given as input the encoding of a BID representation  $J$ , and a binary representation of  $k > 0$ , calculate  $\mu_J(Q)$ , where  $Q[\alpha \theta k]$  is such that  $\text{sk}(Q) = q$ .

The technical problem we address in this work is the complexity of query evaluation. We shall see for the query in Ex. 3, that the query evaluation problem is hard for  $\#\mathcal{P}$ . And moreover, this is the general case for all HAVING queries.

### 3 Preliminaries

We review here some basic facts on semirings (mostly from [12]) and introduce random variables over semirings.

#### 3.1 Background: Random Variables on Semirings

**Definition 5.** A **monoid**  $(S, +, 0)$  is a set  $S$  and  $+$  is an associative binary operation with identity, 0. A **semiring** is a structure  $(S, +, \cdot, 0, 1)$  where  $(S, +, 0)$  is a commutative monoid with identity 0,  $(S, \cdot, 1)$  is a monoid with identity 1,  $\cdot$  distributes over  $+$  and 0 annihilates  $S$ . A commutative semiring is one in which  $(S, \cdot, 1)$  forms a commutative monoid. We abbreviate either structure with the set  $S$  when clear from the context.

We shall consider only commutative semirings in this paper.

*Example 4.* For integer  $k \geq 0$ , let  $\mathbb{Z}_{k+1} = \{0, 1, \dots, k\}$  then for every such  $k$ ,  $(\mathbb{Z}_k, \max, \min, 0, k)$  is a semiring. In particular,  $k = 2$  is the Boolean semiring. Another set of semirings we consider,  $\mathbb{S}_k = (\mathbb{Z}_k, +_k, \cdot_k, 0, 1)$  where  $+_k(x, y) = \min(x + y, k)$  and  $\cdot_k = \min(xy, k)$  where addition and multiplication are in  $\mathbb{Z}$ .

For the rest of this section fix a BID instance  $J$ , and denote  $(\mathcal{W}, \mu) = (\mathcal{W}_J, \mu_J)$  the probability space induced by  $J$  on possible worlds.

**Definition 6.** Given a semiring  $(S, +_S, \cdot_S)$ , an  **$S$ -random variable**,  $r$ , is a function  $r : \mathcal{W} \rightarrow S$ . Given two random variables  $r, s$  then  $r +_S s$  and  $r \cdot_S s$  are random variables defined as  $(r +_S s)(W) = r(W) +_S s(W)$  and  $(r \cdot_S s)(W) = r(W) \cdot_S s(W)$ .

In the sequel, we make use of the following fact:

**Fact 1** The set of  $S$ -random variables on a fixed BID instance  $J$  induces a semiring, denoted  $S^J$ , with the operations in Def. 6.

**Definition 7.** Given a semiring  $S$  and a set of random variables  $R = \{r_1, \dots, r_n\}$  on  $S$ ,  $R$  is **independent** if  $\forall N \subseteq 1, \dots, n$  and any set  $k_1, \dots, k_n \in S$ , we have  $\mu(\bigwedge_{i \in N} r_i = k_i) = \prod_{i \in N} \mu(r_i = k_i)$ . We say that  $R$  is **disjoint** if for any  $i \neq j$  we have  $\mu((r_i \neq 0) \wedge (r_j \neq 0)) = 0$ .

To represent a single random variable we need space as large as  $|W|$ , which is exponential in the size of the  $J$  and thus prohibitive for most applications. However, there exists an alternative representation in terms of *marginal vectors*, which only takes size  $S$ .

**Definition 8.** Given a random variable  $r$  on  $S$ , the **marginal vector** (or simply, the marginal) of  $r$  is denoted  $\mathbf{m}^r$  and is a vector indexed by  $S$  defined by  $\forall s \in S \mu(r = s) = \mathbf{m}^r[s]$ . Given a monoid  $(S, +_S, 0)$ , the **monoid convolution** is a binary operation on marginals denoted  $\ast^{+s}$ , and for any marginals  $\mathbf{m}^r$  and  $\mathbf{m}^s$  is defined by

$$\forall s \in S (\mathbf{m}^r \ast^{+s} \mathbf{m}^s)[s] \stackrel{\text{def}}{=} \sum_{\substack{i+s=j \\ i,j \in S}} \mathbf{m}^r[i] \mathbf{m}^s[j]$$

The disjoint operation for  $(S, 0)$  is denoted  $\mathbf{m}^r \perp \mathbf{m}^s$  and is defined by

$$\text{if } s \neq 0 (\mathbf{m}^r \perp \mathbf{m}^s)[s] \stackrel{\text{def}}{=} \mathbf{m}^r[s] + \mathbf{m}^s[s] \text{ else } (\mathbf{m}^r \perp \mathbf{m}^s)[0] \stackrel{\text{def}}{=} (\mathbf{m}^r[0] + \mathbf{m}^s[0]) - 1$$

The next proposition tells us when the operations defined in the previous definition yield the correct results:

**Proposition 1.** If  $r$  and  $s$  are random variables on the monoid  $(S, +_S, 0)$  with marginal vectors  $\mathbf{m}^r$  and  $\mathbf{m}^s$  then let  $\mathbf{m}^{r+s}$  denote the marginal of  $r +_S s$ . If  $r$  and  $s$  are independent then  $\mathbf{m}^{r+s} = \mathbf{m}^r \ast^{+s} \mathbf{m}^s$ . If  $r$  and  $s$  are disjoint then  $\mathbf{m}^{r+s} = \mathbf{m}^r \perp \mathbf{m}^s$ . Further, the  $n$ -fold convolution can be computed in time  $O(n|S|^2)$  and the  $n$ -fold disjoint operation can be computed in  $O(n|S|)$ .

The importance of this proposition is that if the semiring is small, then each operation can be done efficiently.

*Example 5.* Consider the Boolean semiring, and two random variables  $r$  and  $s$  with marginal probabilities (of truth)  $p_r$  and  $p_s$ . Then  $\mathbf{m}^r = (1 - p_r, p_r)$  and  $\mathbf{m}^s = (1 - p_s, p_s)$ . If  $r$  and  $s$  are independent then, the distribution of  $r \vee s = r +_S s$  which can be computed using  $r \ast^+ s$ . This satisfies  $(r \ast^+ s)[1] = p_r(1 - p_s) + (1 - p_r)p_s + p_r p_s$  or in a more familiar form,  $\mathbf{m}^r \ast^+ \mathbf{m}^s = ((1 - p_r)(1 - p_s), 1 - (1 - p_r)(1 - p_s))$ .

### 3.2 Background: Queries on databases annotated from a semiring

In this section, we review material from [12] on computing queries on databases annotated with semiring elements. A slight twist on prior art is that we allow the query to induce the annotations as a first step, rather than being part of the data.

**Definition 9.** Given a commutative semiring  $S$  and a Boolean conjunctive query  $q = g_1, \dots, g_n$ , an annotation is a set of functions indexed by subgoals, such that for  $i = 1, \dots, n$ ,  $\tau_{g_i}$  is a function on tuples that unify with  $g_i$  to  $S$ . We denote the set of annotation functions with  $\tau$ .

As in [12], we compute the annotation using a modified relational algebra which we define below:

**Definition 10.**

- a **plan**  $P$  is inductively defined as (a) a single subgoal (b)  $\pi_{-x}P_1$  if  $P_1$  is a plan (b)  $P_1 \bowtie P_2$  if  $P_1, P_2$  are plans.
- $\mathbf{var}(P)$ , the variables output by  $P$ , is defined as  $\mathbf{var}(g)$  if  $P = g$ ,  $\mathbf{var}(\pi_{-x}P) = \mathbf{var}(P) - \{x\}$  and  $\mathbf{var}(P_1 \bowtie P_2) = \mathbf{var}(P_1) \cup \mathbf{var}(P_2)$ .
- $\mathbf{goal}(P)$ , the set of subgoals in  $P$ , is defined as (a)  $\mathbf{goal}(g) = \{g\}$ , (b)  $\mathbf{goal}(\pi_{-x}P_1) = \mathbf{goal}(P_1)$  (c)  $\mathbf{goal}(P_1 \bowtie P_2) = \mathbf{goal}(P_1) \cup \mathbf{goal}(P_2)$ .

The **value** of a plan  $P$  on a standard instance  $W$  is a set of tuples with attributes corresponding to the variables in  $\mathbf{var}(P)$  each annotated with a semiring element, denoted  $\omega_P^W(t)$ , which is defined inductively below. Concurrently, we define the support of a tuple  $\mathbf{supp}_{P,V}(t) = \{t' \mid \forall y \in V \ t[y] = t'[y] \wedge \omega_P^W(t') \neq 0\}$  where  $P$  is a plan,  $V$  is a set of variables such that  $V \subseteq \mathbf{var}(P)$ ,  $t$  is a tuple with attributes corresponding to  $V$  and  $t'$  is a tuple with attributes corresponding to  $\mathbf{var}(P)$ .

- If  $P = g$  then if  $t$  unifies with  $g$  then  $\omega_P^W(t) = \tau_g(t)$  else  $\omega_P^W(t) = 0$ .
- If  $P = \pi_{-x}P_1$ , then  $\omega_{\pi_{-x}P_1}^W(t) = \sum_{t' \in \mathbf{supp}_{P_1, \mathbf{var}(P)}^W(t)} \omega_{P_1}^W(t')$ .
- else  $P = P_1 \bowtie P_2$  and for  $i = 1, 2$  let  $t_i$  be  $t$  restricted to  $\mathbf{var}(P_i)$  then  $\omega_{P_1 \bowtie P_2}^W(t) = \omega_{P_1}^W(t_1) \cdot \omega_{P_2}^W(t_2)$

A result of [12] shows that  $\omega_P^W$  is independent of the choice of plan  $P$ , which justifies the notation  $s_{\tau, W, q}$ , the value of a conjunctive query  $q$  on a deterministic instance  $W$  under annotation  $\tau$  defined as  $s_{\tau, W, q} \stackrel{\text{def}}{=} \omega_P^W()$  where  $P$  is any plan for  $q$  where and  $\omega_P^W$  is applied to the empty tuple.

## 4 Approaches for HAVING

In this section, we define the  $\alpha$ -safe HAVING queries for  $\alpha \in \{\text{MIN}, \text{MAX}, \text{COUNT}\}$  in Sec. 4.3, for  $\alpha = \text{COUNT}(\text{DISTINCT})$  in Sec. 4.4 and  $\alpha \in \{\text{AVG}, \text{SUM}\}$  in Sec. 4.5.

### 4.1 Aggregates and semirings

We explain the details of computing HAVING queries using semirings on deterministic databases, which immediately generalizes to probabilistic databases. Since HAVING queries are Boolean, we use a function  $\rho$ , the **recovery function**, which maps semiring values to true if that value would satisfy the HAVING query. Fig. 3 lists the (commutative) semirings for the aggregates in this paper, their annotations  $\tau$  and a Boolean recovery function  $\rho$ . EXIST is similar to the safe plan algebra of [7, 8, 21].

| HAVING Predicate       | Semiring                     | Annotation $\tau_{g^*}(t)$                         | Recovery $\rho(s)$                               |
|------------------------|------------------------------|----------------------------------------------------|--------------------------------------------------|
| EXISTS                 | $(\mathbb{Z}_2, \max, \min)$ | 1                                                  | $s = 1$                                          |
| MIN(y) $\{<, \leq\} k$ | $(\mathbb{Z}_3, \max, \min)$ | if $t \theta k$ then 2 else 1                      | $s = 2$                                          |
| MIN(y) $\{>, \geq\} k$ | $(\mathbb{Z}_3, \max, \min)$ | if $t \theta k$ then 1 else 2                      | $s = 1$                                          |
| MIN(y) $\{=, \neq\} k$ | $(\mathbb{Z}_4, \max, \min)$ | if $t < k$ then 3 else<br>if $t = k$ then 2 else 1 | if $=$ then $s = 2$<br>if $\neq$ then $s \neq 2$ |
| COUNT(*) $\theta k$    | $\mathbb{S}_{k+1}$           | 1                                                  | $(s \neq 0) \wedge (s \theta k)$                 |
| SUM(y) $\theta k$      | $\mathbb{S}_{k+1}$           | $t[y]$                                             | $(s \neq 0) \wedge (s \theta k)$                 |

**Fig. 3.** Semirings for the operators MIN, COUNT and SUM. Let  $g^*$  be the lowest indexed subgoal such that contains  $y$ . For all  $g \neq g^*$ ,  $\forall t$ ,  $\tau_g(t)$  equals the multiplicative identity of the semiring. Let  $\mathbb{Z}_{k+1} = \{0, 1, \dots, k\}$  and  $+_k(x, y) \stackrel{\text{def}}{=} \min(x + y, k)$  and  $\cdot_k \stackrel{\text{def}}{=} \min(xy, k)$ , where  $x, y \in \mathbb{Z}$ . Let  $\mathbb{S}_k \stackrel{\text{def}}{=} (\mathbb{Z}_{k+1}, +_k, \cdot_k, 0, 1)$ . MAX and MIN are symmetric. AVG and COUNT(DISTINCT) are omitted.

*Example 6.* Consider the query  $Q[\text{MIN}(y) \geq 10] :- R(y)$  where  $R = \{t_1, \dots, t_n\}$ . Fig. 3 tells us to use the semiring  $(\mathbb{Z}_3, \max, \min)$ . We first apply  $\tau$ :  $\tau(t_i) = 1$  represents that  $t_i[y] > 10$  while  $\tau(t_i) = 2$  represents that  $t_i[y] \leq 10$ . Let  $s = \sum_{i=1, \dots, m}^S \tau(t_i) = \max_{i=1, \dots, m} \tau(t_i)$ .  $\rho(s)$  is satisfied only when  $s$  is 1, i.e. all  $t_i[y]$  are greater than 10.

More generally, we have the following proposition:

**Proposition 2.** *Given a HAVING query  $Q$ , let  $q = \text{sk}(Q)$  and  $S, \rho$  and  $\tau$  be chosen as in Fig. 3, then for any deterministic instance  $W$  and  $s_{\tau, W, q}$  (Sec. 3.2):*

$$W \models Q \iff \rho(s_{\tau, W, q})$$

In probabilistic databases, we want to compute the random variable  $s_{\tau, q}$  defined as  $s_{\tau, q}(W) \stackrel{\text{def}}{=} s_{\tau, W, q}$ . A simple corollary of Prop. 2 is the following generalization to probabilistic databases:

**Corollary 1.** *Given  $Q$ , let  $q, S, \rho$  and  $\tau$  be as in Prop. 2 then for any BID instance  $J$  we have the following equalities:*

$$\mu_J(Q) = \mu_J(\rho(s_{\tau, W, q})) = \sum_{k: \rho(k)} \mathbf{m}^{s_{\tau, q}}[k]$$

Cor. 1 tells us that examining the entries of the marginal vector at index  $s$  such that  $\rho(s)$  is true is sufficient to answer  $Q$ . Hence our goal is to compute  $\mathbf{m}^{s_{\tau, q}}$ .

## 4.2 Computing safely in semirings

We now extend safe plans to compute a marginal vector instead of a Boolean value. Specifically, we compute  $\mathbf{m}^{s_{\tau, q}}$ , the marginal vector for  $s_{\tau, q}$  using the operations defined in Sec. 3.1.

**Definition 11.** An extensional plan for a Boolean conjunctive query  $q$  is a subgoal  $g$  and if  $P_1, P_2$  are extensional plans then so are  $\pi_{-x}^I P_1$ ,  $\pi_{-x}^D P_1$  and  $P_1 \bowtie P_2$ . An extensional plan  $P$  is **safe** if  $P_1$  and  $P_2$  are safe then

- $P = g$  is safe
- $P = \pi_{-x}^I P_1$  is safe if  $x \in \mathbf{var}(P_1)$  and  $\forall g \in \mathbf{goal}(P_1)$  then  $x \in \mathbf{key}(g)$
- $P = \pi_{-x}^D P_1$  is safe if  $x \in \mathbf{var}(P_1)$  and  $\exists g \in \mathbf{goal}(P_1)$ ,  $\mathbf{key}(g) \subseteq \mathbf{var}(P)$ ,  $x \in \mathbf{var}(g)$ .
- $P = P_1 \bowtie P_2$  is safe if  $\mathbf{goal}(P_1) \cap \mathbf{goal}(P_2) = \emptyset$ 
  - and for  $i = 1, 2$   $\mathbf{var}(\mathbf{goal}(P_1)) \cap \mathbf{var}(\mathbf{goal}(P_2)) \subseteq \mathbf{var}(P_i)$

An extensional plan  $P$  is a safe plan for  $q$  if  $P$  is safe and  $\mathbf{goal}(P) = q$ .

**Proposition 3.** If  $P$  is a safe plan for  $q$ , then for  $x \in \mathbf{var}(q)$  there is exactly one of  $\pi_{-x}^I$  or  $\pi_{-x}^D$  in  $P$ .

At least one of the two projections must be present, because we must remove the variable  $x$  ( $q$  is boolean). If there were more than one in the plan, then they cannot be descendants of each other because  $x \notin \mathbf{var}(P_1)$  for the ancestor and they cannot be joined afterward because of the join condition for  $i = 1, 2$   $\mathbf{var}(\mathbf{goal}(P_1)) \cap \mathbf{var}(\mathbf{goal}(P_2)) \subseteq \mathbf{var}(P_i)$

**Definition 12.** Given a BID instance  $J$ . Let  $P$  be an safe plan, then denote the **extensional value** of  $P$  in  $S$  on  $J$  as  $\hat{\omega}_{P,S}^J$  which is a marginal vector on  $S$  defined inductively:

- $\hat{\omega}_{g,S}^J(t) = \hat{\tau}_g(t)$  where  $\hat{\tau}_g(t)$  is the marginal vector  $\mathbf{m}^t$  given by  $\mathbf{m}^t[0] = 1 - t[P]$  and  $\mathbf{m}^t[\tau_g(t)] = t[P]$ , i.e. the (probabilistic) image of  $\tau$ .
- $\hat{\omega}_{\pi_{-x}^I P_1, S}^J(t) = \ast_{t' \in \mathbf{supp}_{P_1, \mathbf{var}(P_1)}(t)}^{+S} \hat{\omega}_{P_1, S}^J(t)$ .
- $\hat{\omega}_{\pi_{-x}^D P_1, S}^J(t) = \perp_{t' \in \mathbf{supp}_{P_1, \mathbf{var}(P_1)}(t)} \hat{\omega}_{P_1, S}^J(t)$ .
- $\hat{\omega}_{P_1 \bowtie P_2}^J(t) = \hat{\omega}_{P_1, S}^J(t_1) \ast \hat{\omega}_{P_2, S}^J(t_2)$  where for  $i = 1, 2$   $t_i$  is  $t$  restricted to  $\mathbf{var}(P_i)$

The next lemma uses the observation that an operator in a safe plan and the operator used to compute the value have the same correlation assumptions. For example,  $\pi^I$  assumes independence, which is required for  $\ast^+$ .

**Lemma 1.** If  $P$  is a safe plan for a Boolean query  $q$  then for any  $s_i \in S$  on any BID instance  $J$ , we have  $\hat{\omega}_P^J(s_i) = \mu_J(s_{\tau, q} = s_i)$ .

*Remark 2.* A safe plan is not necessarily efficient for any  $S$  (Def. 4). In particular, the operations in a safe plan on  $S$  take time polynomial in  $|S|$ . Thus, if the size of  $S$  grows super-polynomially in  $|J|$ , the size of the BID instance, the plan will not be efficient.

### 4.3 MIN, MAX and COUNT-safe

We can now formalize the class of queries which are efficient for MIN, MAX and COUNT.

**Definition 13.** If  $\alpha \in \{\text{MIN}, \text{MAX}, \text{COUNT}\}$  and  $Q[\alpha(t) \theta k]$  is a HAVING query, then  $Q$  is  $\alpha$ -safe if the skeleton of  $Q$  is safe.

**Theorem 1.** *If  $Q[\alpha(y) \theta k]$  is a **HAVING** query for  $\alpha \in \{\text{MIN}, \text{MAX}, \text{COUNT}\}$  and  $Q$  is  $\alpha$ -safe then the exact evaluation problem for  $Q$  is in polynomial time.*

Correctness is straightforward from Lem. 1. Efficiency follows because the semiring is of polynomial size. Hence the translation and each operation in the evaluation is of polynomial size for each aggregate in the theorem. In particular,  $S$  is constant for MIN and MAX and upper bounded by  $n$  for COUNT.

*Remark 3.* We remark that for SUM, we can only guarantee that  $|S| = O(k)$ , which implies a running time of  $O(kn^{|Q|})$ , which is not efficient (Def. 4).

The results of [7, 8, 21] show that either a conjunctive query without self-joins has a safe plan or it is  $\#P$ -hard. It is not hard to see that each aggregate has at least one test so that a **HAVING** query  $Q$  is satisfied only when the skeleton of  $Q$  is satisfied. It is then straightforward to extend to all predicate tests. Formally, we have:

**Theorem 2.** *If  $\alpha \in \{\text{MIN}, \text{MAX}, \text{COUNT}\}$  and  $Q[\alpha(y) \theta k]$  does not contain self-joins, if  $Q$  is  $\alpha$ -safe then  $Q$  has data complexity in  $\mathcal{P}$  else  $Q$  has  $\#P$ -hard data complexity. Further, we can find an  $\alpha$ -safe plan in  $\mathcal{P}$ .*

The algorithm to find a safe plan is identical to [8, 21].

#### 4.4 COUNT(DISTINCT)-safe queries

Intuitively, we compute COUNT(DISTINCT) in two stages: proceeding bottom-up, we first compute the probability a  $y$  value appears (i.e. DISTINCT), we then count the number of distinct values (i.e. COUNT) using the techniques of the previous section. However, one caveat is that the representation is lossy: We do not know *which* values are present, only the distribution of their count. This implies that not all operations on these lossy marginal vectors are correct, which restricts the class of allowable plans:

**Definition 14.** *A query  $Q[\text{COUNT}(\text{DISTINCT } y) \theta k]$  is **COUNT(DISTINCT)-safe** if there is a safe plan  $P$  for the skeleton of  $Q$  such that  $\pi_{-y}^L$  or  $\pi_{-y}^D$  is in  $P$  and no proper ancestor is  $\pi_{-x}^L$  for any  $x$ .*

*Example 7.* Fix a BID instance  $J$ . Consider  $Q[\text{COUNT}(\text{DISTINCT } y) > 2] :- R(y, x), S(y)$ , a safe plan for the skeleton of  $Q$  is  $P = \pi_{-y}^L((\pi_{-x}^L R(y, x)) \bowtie S(y))$ . For the subquery  $P_1 = (\pi_{-x}^L R(y, x)) \bowtie S(y)$ , calculate the probability that each value EXISTS in the subplan  $P_1$ , i.e. for each  $t$ ,  $\hat{\omega}_{P, \text{EXISTS}}^J(t)$ . Intuitively, all tuples returned by the plan are independent because  $\pi_{-y}^L$  is correct, and all  $y$  values are trivially distinct.

Intuitively, we map each EXISTS marginal vector to a vector suitable for computing COUNT, i.e. a vector in  $\mathbb{Z}_k$  where  $k = 2$ , in this problem. In other words,  $(1 - p, p) = \hat{\omega}_{P, \text{EXISTS}}^J(t) = \mathbf{m}'$  is mapped to  $\hat{\tau}(\mathbf{m}') = (1 - p, p, 0)$ . Thus, the correct distribution is given by  $*_{t' \in \text{supp}_P(t)} \hat{\tau}(t')$ . To compute the final result, use the recovery function,  $\rho$  defined by  $\rho(s) \iff s > 2$

The proof of the following theorem is a generalization of Ex. 7, whose proof is in the full paper [23]:

**Theorem 3.** *If  $Q$  is COUNT(DISTINCT)-safe then its evaluation problem is in  $\mathcal{P}$ .*

**Complexity** We now establish that for COUNT(DISTINCT) queries without self-joins, COUNT(DISTINCT)-safe captures efficient computation. We first show the canonical hard patterns for COUNT(DISTINCT) and then extend this to show that the evaluation of any COUNT(DISTINCT) query without self-joins that is not COUNT(DISTINCT)-safe can be reduced to one of these hard patterns.

**Proposition 4.** *The following HAVING queries are #P-hard for  $i \geq 1$ :*

$$Q_1[\text{COUNT}(\text{DISTINCT } y) \theta k] :- R^p(x), S(x, y) \text{ and } Q_{2,i}[\text{COUNT}(\text{DISTINCT } y) \theta k] :- R_1^p(x; y), \dots, R_i^p(x; y)$$

*Proof (Sketch).* To see that  $Q_1$  is hard, we reduce from counting the number of independent sets in a graph  $(V, E)$ . We let  $k$  be the number of edges, intuitively  $Q$  is satisfied only when all edges are present. For each node  $u \in V$ , create a tuple  $R(u)$  with probability 0.5. For edge  $e = (u, v)$  create two tuples  $S(e, u), S(e, v)$ , each with probability 1. For any set  $V' \subseteq V$ , let  $W_{V'}$  denote the world corresponding where the tuples corresponding to  $V'$  are present. For any subset of nodes,  $N$ , we show that if  $N$  is an independent set if and only if  $W_{V-N}$  satisfies  $Q_1$ . Since  $f(N) = V - N$  is one-to-one, the number of possible worlds that satisfy  $Q_1$  are exactly the number of independent sets. If  $N$  is independent, then for any edge  $(u, v)$ , it must be the case that at least one of  $u$  or  $v$  is in  $V - N$ , thus every edge is present and  $Q$  is satisfied. If  $N$  is not independent, then there must be some edge  $(u, v)$  such that  $u, v \in N$ , hence neither of  $u, v$  is in  $V - N$ . Since this edge is missing,  $Q_1$  cannot be satisfied. The hardness of  $Q_2$  is based on a reduction from counting the set covers of a fixed size and is in the full paper [23].

There is some work in generalizing the previous theorem, we show in the full paper [23]:

**Theorem 4.** *If  $Q$  is not COUNT(DISTINCT)-safe and does not contain self-joins, then  $Q$  has #P-hard data complexity.*

*Proof (Sketch).* We sketch the proof in the simpler case when only tuple independent probabilistic tables are used in  $Q$ . Assume the theorem fails, let  $Q$  be the minimal counter example in terms of subgoals; this implies we may assume that  $Q$  is connected and the skeleton of  $Q$  is safe. Since there is no safe plan projecting on  $y$  and only independent projects are possible, the only condition that can fail is that some subgoal does not contain  $y$ . Thus, there are at least two subgoals  $R(x)$  and  $S(zy)$  such that  $y \notin x \cup z$  and  $x \cap z \neq \emptyset$ . Given a graph  $(V, E)$ , we then construct a BID instance  $J$  exactly as in the proof of Prop. 4. Only the  $R$  relation is required to have probabilistic tuples, all others can set their probabilities to 1.

Extending to tuple-disjoint databases requires slightly more work, because adding multiple tuples with probability 1 may violate a possible world key constraint. The proof appears in the full paper [23]. It is straightforward to decide if a plan is COUNT(DISTINCT)-safe, the safe plan algorithm of [8, 21] simply tries only disjoint projects and joins until it is able to project away  $y$  or it fails.

#### 4.5 SUM-safe and AVG-safe queries

To find SUM- and AVG-safe queries, we further restrict the class of allowable plans. Intuitively, if each value for  $y$  is present disjointly, then computing the multiplicity of each value is sufficient to compute the query, which we accomplish using the COUNT algebra of Sec. 4.3. Since computing SUM and AVG for a HAVING query with a single tuple independent table is already  $\#P$ -Hard, our safe plans we must not contain an independent project for  $y$  ( $\pi_{-y}^I$ ).

**Definition 15.** A HAVING query  $Q[\alpha(y) \theta k]$  for  $\alpha \in \{SUM, AVG\}$  is  $\alpha$ -safe, if there is a safe plan  $P$  for the skeleton of  $Q$  such that  $\pi_{-y}^D$  in  $P$  and no proper ancestor of  $\pi_{-y}^D$  is  $\pi_{-x}^I$  for any  $x$ .

**Theorem 5.** If  $Q[\alpha(y) \theta k]$  for  $\alpha \in \{SUM, AVG\}$  is  $\alpha$ -safe, then  $Q$ 's evaluation problem is in  $\mathcal{P}$ .

*Proof (Sketch).* Since  $Q$  is  $\alpha$ -safe, there is a plan  $P$  satisfying Def. 15. We may assume without loss of generality that  $P = \pi_{-y}^D(P_1)$ . We compute  $P_1$  using the algebra for COUNT,  $\hat{\omega}_{\text{COUNT}, P_1}^J(t)$ , which is correct because  $P_1$  is safe. We then have the following equation which can be computed efficiently:

$$\mu(Q) = \sum_{t \in \text{supp}_{P_1}()} \sum_{s: s * t[y] \theta k} \hat{\omega}_{P_1, \text{COUNT}}^J[s]$$

*Example 8.* Consider  $Q[\text{SUM}(y) > 10] :- R('a'; y), S(y, u)$ . This query is SUM-safe, with plan  $\pi_{-y}^D(R('a'; y) \bowtie \pi_{-u}^I S(y, u))$ .

**Complexity** We show that if a HAVING query without self-joins is not SUM-safe then, it has  $\#P$ -data complexity. AVG follows by essentially the same construction.

**Proposition 5.** If  $\alpha \in \{SUM, AVG\}$  and  $\theta \in \{=, <, =, >, \geq\}$  then  $Q[\alpha(y) \theta k] :- R^p(y)$  has  $\#P$ -data complexity.

*Proof (Sketch).* Consider when  $\theta$  is  $=$ . An instance of  $\#SUBSET$ -SUM is a set of integers  $x_1, \dots, x_n$  and our goal is to count the number of subsets  $\emptyset \neq S \subseteq 1, \dots, n$  such that  $\sum_{s \in S} x_s = B$ . We create the representation with schema  $R(X; ; P)$  satisfying  $R = \{(x_1; 0.5), \dots, (x_n; 0.5)\}$ , i.e. each tuple present with probability 0.5. Thus,  $\mu(Q) * 2^n$  is number of such  $S$ . Showing hardness for other aggregate tests is straightforward.

**Theorem 6.** If  $\alpha \in \{SUM, AVG\}$  and  $\theta \in \{=, \neq, \leq, <, >, \geq\}$  then let  $Q[\alpha(y) \theta k]$  be a HAVING query if  $Q$  does not contain self-joins and is not  $\alpha$ -safe, then  $Q$  has  $\#P$ -data complexity. Further, there is an algorithm to decide if  $Q$  is  $\alpha$ -safe in  $\mathcal{P}$ .

If  $Q$ 's skeleton is unsafe, then it is straightforward that  $Q$  has  $\#P$ -hard data complexity. So we may assume that  $Q$  is safe, but not SUM-safe. Given any set of constants  $y_1, \dots, y_n$ , let  $q = \text{sk}(Q)$  and for  $i = 1, \dots, n$ ,  $q_i$  be  $q[y \rightarrow y_i]$ . We construct an instance  $J$  such that the  $q_i$  are satisfied independently with multiplicity 1<sup>3</sup>. This allows us to construct the reduction above. A formal proof is in the full paper [23].

<sup>3</sup> By multiplicity, we mean that for any  $W \in \mathcal{W}_J$  and any  $q_i$  there is at most one valuation  $v$  such that  $\text{im}(v) \subseteq W$

## 5 Related Work

Probabilistic relational databases have been discussed by Barbara et. al [3] and more recently Dalvi and Suciu [7], Ré et. al [22], Sen et. al [27] and Widom [29], though all omit HAVING style aggregation. Cheng et al. [6] and Desphande et. al [9] consider probabilistic databases resulting from sensor networks and handle continuous distributions with more general correlations, while we handle only the discrete case. In their settings, aggregate queries are effectively value aggregates over a single relation.

In the OLAP setting [4] and streaming setting [16] give efficient algorithms for *value aggregation* in a model which is equivalent to the single table model and focuses on scaling such computation (e.g. using streaming techniques). In contrast, computing the AVG for predicate aggregates on a single table is already  $\#P$ -Hard. Ross et al. [25] describe an approach to computing aggregates on a probabilistic database, by computing bounding intervals (e.g. the AVG is between [5600, 5700]). For more aggregate functions than we discuss, they show computing bounding intervals exactly is  $NP$ -Hard but do not offer any results on the boundary of hardness.

A closely related work is Arenas et. al, [2] which considers the complexity of aggregate queries, similar to HAVING queries, over data which violates functional dependencies. Their semantic is greatest lower bound or least upper bound on the set of all minimal repairs, i.e. not probabilistic. They consider multiple predicates, which we do not. In this paper, we deal with more general types of value inconsistency.

## 6 Conclusion

In this paper we have examined the complexity of evaluating positive conjunctive queries with predicate aggregates over probabilistic databases. For each aggregate, we discussed a novel method based on computing the distribution of elements in a semiring to evaluate such queries. We proved that for conjunctive queries without self-joins our methods are optimal.

**Acknowledgements** This work was partially supported by NSF Grants IIS-0428168 and IIS-0454425.

## References

1. S. Abiteboul, R. Hull, and V. Vianu. *Foundations of Databases*. Addison-Wesley, 1995.
2. M. Arenas, L. Bertossi, J. Chomicki, X. He, V. Raghavan, and J. Spinrad. Scalar aggregation in inconsistent databases. *Theoretical Computer Science.*, 2003.
3. D. Barbara, H. Garcia-Molina, and D. Porter. The management of probabilistic data. *IEEE Trans. Knowl. Data Eng.*, 4(5):487–502, 1992.
4. D. Burdick, P. M. Deshpande, T. S. Jayram, R. Ramakrishnan, and S. Vaithyanathan. Olap over uncertain and imprecise data. *VLDB J.*, 16(1):123–144, 2007.
5. M.J. Cafarella, C. Ré, D. Suciu, and O. Etzioni. Structured querying of web text data: A technical challenge. In *CIDR*, pages 225–234. [www.crdrrdb.org](http://www.crdrrdb.org), 2007.
6. R. Cheng, D. Kalashnikov, and S. Prabhakar. Evaluating probabilistic queries over imprecise data. In *Proceedings of ACM SIGMOD Conference*, 2003.

7. N. Dalvi and D. Suciu. Efficient query evaluation on probabilistic databases. In *VLDB*, Toronto, Canada, 2004.
8. N. Dalvi and D. Suciu. Management of probabilistic data: Foundations and challenges. In *PODS*, pages 1–12, 2007.
9. A. Deshpande, C. Guestrin, S. Madden, J. Hellerstein, and W. Hong. Model-driven data acquisition in sensor networks, 2004.
10. A. Fuxman and R. J. Miller. First-order query rewriting for inconsistent databases. In *ICDT*, pages 337–351, 2005.
11. E. Gradel, Yu. Gurevich, and C. Hirsch. The complexity of query reliability. In *Symposium on Principles of Database Systems*, pages 227–234, 1998.
12. T. Green, G. Karvounarakis, and V. Tannen. Provenance semirings. In *PODS*, 2007.
13. T.J. Green and Val Tannen. Models for incomplete and probabilistic information. *IEEE Data Engineering Bulletin*, 29, 2006.
14. R. Gupta and S. Sarawagi. Curating probabilistic databases from information extraction models. In *Proc. of the 32nd Int'l Conference on Very Large Databases (VLDB)*, 2006.
15. M.A. Hernandez and S.J. Stolfo. The merge/purge problem for large databases. In *SIGMOD Conference*, pages 127–138, 1995.
16. T.S. Jayram, S. Kale, and E. Vee. Efficient aggregation algorithms for probabilistic data. In *SODA*, 2007.
17. T.S. Jayram, R. Krishnamurthy, S. Raghavan, S. Vaithyanathan, and H. Zhu. Avatar information extraction system. *IEEE Data Engineering Bulletin*, 29(1), 2006.
18. L. Lakshmanan, N. Leone, R. Ross, and V.S. Subrahmanian. Proview: A flexible probabilistic database system. *ACM Trans. Database Syst.*, 22(3), 1997.
19. I. Mansuri and S. Sarawagi. A system for integrating unstructured data into relational databases. In *Proc. of the 22nd IEEE Int'l Conference on Data Engineering (ICDE)*, 2006.
20. A. Parag, O. Benjelloun, A.D. Sarma, C. Hayworth, S. Nabar, T. Sugihara, and J. Widom. Trio: A system for data uncertainty and lineage. In *VLDB*, 2006.
21. C. Ré, N. Dalvi, , and D. Suciu. Query evaluation on probabilistic databases. *IEEE Data Engineering Bulletin*, 29(1):25–31, 2006.
22. C. Ré, N. Dalvi, and D. Suciu. Efficient top-k query evaluation on probabilistic data. In *Proceedings of ICDE*, 2007.
23. C. Ré and D. Suciu. Efficient evaluation of having queries on a probabilistic database. Technical Report TR2007-06-01, University of Washington, Seattle, Washington, June 2007.
24. C. Ré and D. Suciu. Materialized views in probabilistic databases for information exchange and query optimization. In *VLDB*, 2007.
25. Robert Ross, V. S. Subrahmanian, and John Grant. Aggregate operators in probabilistic databases. *J. ACM*, 52(1):54–101, 2005.
26. A.D. Sarma, O. Benjelloun, A.Y. Halevy, and J. Widom. Working models for uncertain data. In Ling Liu, Andreas Reuter, Kyu-Young Whang, and Jianjun Zhang, editors, *ICDE*, page 7. IEEE Computer Society, 2006.
27. P. Sen and A. Deshpande. Representing and querying correlated tuples in probabilistic databases. In *Proceedings of ICDE*, 2007.
28. L.G. Valiant. The complexity of enumeration and reliability problems. *SIAM J. Comput.*, 8(3):410–421, 1979.
29. J. Widom. Trio: A system for integrated management of data, accuracy, and lineage. In *CIDR*, pages 262–276, 2005.
30. W.E. Winkler. Improved decision rules in the fellegi-sunter model of record linkage. Technical report, Statistical Research Division, U.S. Census Bureau, Washington, DC, 1993.
31. W.E. Winkler. The state of record linkage and current research problems. Technical report, Statistical Research Division, U.S. Bureau of the Census, 1999.