

Acyclic Join Sampling Under Selections: Dichotomy, Union Sampling, and Enumeration

Jinchao Huang ✉ 

The Chinese University of Hong Kong, Hong Kong

Yufei Tao ✉ 

The Chinese University of Hong Kong, Hong Kong

Sibo Wang ✉ 

The Chinese University of Hong Kong, Hong Kong

Abstract

Previous research on *join sampling* has focused on joins without selection conditions, even though such conditions are prevalent in everyday queries in database systems. Motivated by this, we undertake a systematic investigation on the complexity of sampling from the result of an acyclic join under equality conditions given only at runtime. When conditions are conjunctive, the goal is to understand when it is possible to precompute a *feasible* structure that uses $\tilde{O}(\text{IN})$ space and supports sampling in $\tilde{O}(1)$ time, where IN is the input size. We present a dichotomy to characterize (subject to a widely-accepted conjecture) the existence of such structures based on the conditions supplied and, in every feasible scenario, give an optimal structure of $O(\text{IN})$ space and $O(1)$ sample time. We then extend our investigation to conditions expressed in disjunctive normal form, where the core challenge reduces to the fundamental *set union sampling* problem. We overcome the challenge with an optimal algorithm and utilize it to develop optimal sampling structures. Our findings also lead to new results on the closely-related *random enumeration* problem.

2012 ACM Subject Classification Theory of computation → Data structures and algorithms for data management

Keywords and phrases Conjunctive Queries, Acyclic Joins, Sampling, Lower Bounds

Digital Object Identifier 10.4230/LIPIcs.ICDT.2026.9

Funding This work was supported in part by GRF Projects 14217322 and 14222822.

1 Introduction

Join evaluation on massive datasets often incurs prohibitive computation. A primary cause of this phenomenon is the sheer volume of the join result – whose size grows rapidly as the number of participating relations increases – because reporting each result tuple requires at least constant time. However, in many scenarios (ranging from approximate query processing and query optimization to interactive data exploration and machine learning), the goal is not to enumerate the entire join result, but rather to obtain a small representative sample. Such samples can be used to approximate aggregates, estimate selectivities, support interactive analysis, or train a learning model without materializing the full join. These practical needs have motivated the development of numerous join sampling algorithms in the theory community; see [2, 8, 13, 15, 16, 19, 28, 32, 35] and the references therein.

Previous research has focused on sampling from joins without selection conditions. This stands in stark contrast to reality, where queries nearly always include such conditions, typically in the form of equality predicates like $(A = a \wedge B = b) \vee (C = c)$. Here, the values a , b , and c are unknown in advance and are supplied by a query only at runtime. Essential in narrowing down the data of interest, these predicates prompt the question of how join sampling can remain effective in the presence of *conjunctive* and *disjunctive* selection conditions. This paper addresses the question in a systematic manner.



© Jinchao Huang, Yufei Tao, and Sibow Wang;
licensed under Creative Commons License CC-BY 4.0

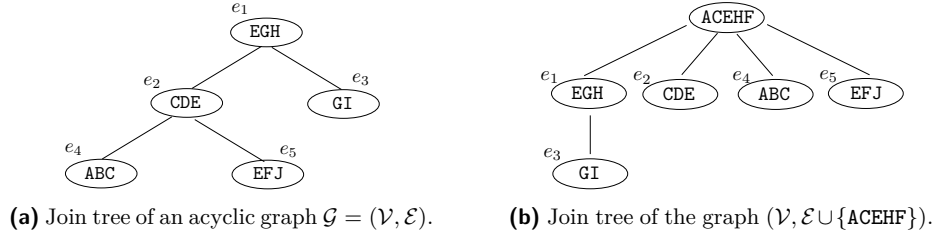
29th International Conference on Database Theory (ICDT 2026).

Editors: Balder ten Cate and Maurice Funk; Article No. 9; pp. 9:1–9:20

Leibniz International Proceedings in Informatics



LIPICs Schloss Dagstuhl – Leibniz-Zentrum für Informatik, Dagstuhl Publishing, Germany



■ **Figure 1** Acyclicity and connectivity.

Mathematical Conventions and Computation Model. The notation $\mathbb{N}$ represents the set of integers. Given an integer $x \geq 1$, we define $[x] = \{1, 2, \dots, x\}$. All logarithms have base 2 by default. Our analysis assumes the standard word-RAM model [17] where a word has length $\Theta(\log \text{IN})$ with IN being the input size. Given an integer x representable using a word, we assume that a uniformly random number can be drawn from $[x]$ in constant time.

1.1 Problem Definitions

Let $\mathbf{att}$ be a set whose elements are called *attributes*. Given a subset $U \subseteq \mathbf{att}$, a *tuple* over U is a function $\mathbf{u} : U \rightarrow \mathbb{N}$. For each attribute $X \in U$, we call $\mathbf{u}(X)$ “the value of $\mathbf{u}$ under X ” or simply the “ X -value” of $\mathbf{u}$ – such a value is assumed to fit in a word. Given a subset U' of U , we use $\mathbf{u}[U']$ (note: square bracket here) to represent the tuple $\mathbf{v}$ over U' satisfying $\mathbf{v}(X) = \mathbf{u}(X)$ for all $X \in U'$; the tuple $\mathbf{u}[U']$ is called the *projection* of $\mathbf{u}$ onto U' . We define a *relation* as a set R of tuples over an identical set U of attributes, where U is called the *schema* of R , written as $\text{schema}(R) = U$.

Let $\mathcal{V}$ be a finite subset of $\mathbf{att}$ and $\mathcal{E}$ be a subset of $2^{\mathcal{V}}$. We call the hypergraph $\mathcal{G} = (\mathcal{V}, \mathcal{E})$ a *schema graph* if every attribute of $\mathcal{V}$ appears in at least one element – a hyperedge – of $\mathcal{E}$. The hypergraph is *acyclic* [1, 33] if there exists a tree $\mathcal{T}$ having two properties:

- [the one-one property] every node of $\mathcal{T}$ corresponds to one distinct hyperedge in $\mathcal{E}$;
- [the connectedness property] for each attribute $X \in \mathcal{V}$, the nodes in $\mathcal{T}$ (a.k.a. hyperedges in $\mathcal{E}$) containing X form a connected subtree in $\mathcal{T}$.

We refer to $\mathcal{T}$ as a *join tree* of $\mathcal{G}$.

► **Example 1.** Consider the hypergraph $\mathcal{G} = (\mathcal{V}, \mathcal{E})$ where $\mathcal{V} = \{\text{A}, \text{B}, \dots, \text{J}\}$, and $\mathcal{E}$ has 5 hyperedges: $e_1 = \text{EGH}$ (shortform for set $\{\text{E}, \text{G}, \text{H}\}$), $e_2 = \text{CDE}$, $e_3 = \text{GI}$, $e_4 = \text{ABC}$, and $e_5 = \text{EFJ}$. The hypergraph is acyclic, with a join tree $\mathcal{T}$ given in Figure 1(a). The hyperedges containing attribute, for instance, E are e_1, e_2 , and e_5 , and they form a connected subtree of $\mathcal{T}$. As another example, the hyperedges containing G are e_1 and e_3 , which again are connected. ◻

Fix a schema graph $\mathcal{G} = (\mathcal{V}, \mathcal{E})$. A *join instance* (or simply, a *join*) of $\mathcal{G}$ is a set $\mathcal{Q}$ of relations such that (i) $|\mathcal{Q}| = |\mathcal{E}|$, and (ii) for each hyperedge $e \in \mathcal{E}$, there is a unique relation $R_e \in \mathcal{Q}$ with $\text{schema}(R_e) = e$. We call $\mathcal{G}$ the “schema graph of $\mathcal{Q}$ ”. The *result* of the join is a relation over $\mathcal{V}$ defined as:

$$\text{Join}(\mathcal{Q}) = \{\text{tuple } \mathbf{u} \text{ over } \mathcal{V} \mid \forall R \in \mathcal{Q}, \mathbf{u}[\text{schema}(R)] \in R\}. \quad (1)$$

The *input size* of $\mathcal{Q}$ is defined as $\text{IN} = \sum_{R \in \mathcal{Q}} |R|$. If $\mathcal{G}$ is acyclic, we refer to $\mathcal{Q}$ as an *acyclic join*; otherwise, it is a *cyclic join*.

If the relations in $\mathcal{Q}$ can be listed as $R_1, R_2, \dots, R_{|\mathcal{Q}|}$, we may also represent $\text{Join}(\mathcal{Q})$ as $R_1 \bowtie R_2 \bowtie \dots \bowtie R_{|\mathcal{Q}|}$. If a relation, say, R_1 consists of a single tuple $\mathbf{u}$, we may replace it with $\mathbf{u}$ in a join expression, e.g., $\mathbf{u} \bowtie R_2 \bowtie \dots \bowtie R_{|\mathcal{Q}|}$.

($\mathcal{G}, \mathcal{Z}$)-Sampling. Fix a schema graph $\mathcal{G} = (\mathcal{V}, \mathcal{E})$ whose size is assumed to be a constant, together with a subset $\mathcal{Z} \subseteq \mathcal{V}$. Let $\mathcal{Q}$ be a join instance of $\mathcal{G}$. Given a tuple $\mathbf{z}$ over $\mathcal{Z}$, define

$$\sigma_{\mathbf{z}}(\text{Join}(\mathcal{Q})) = \{\mathbf{u} \in \text{Join}(\mathcal{Q}) \mid \forall Z \in \mathcal{Z}, \mathbf{u}(Z) = \mathbf{z}(Z)\} \quad (2)$$

namely, $\sigma_{\mathbf{z}}(\text{Join}(\mathcal{Q}))$ comprises the tuples $\mathbf{u}$ in the join result satisfying the conjunctive condition $\bigwedge_{Z \in \mathcal{Z}} \mathbf{u}(Z) = \mathbf{z}(Z)$. Given a tuple $\mathbf{z}$ over $\mathcal{Z}$, a $(\mathcal{G}, \mathcal{Z})$ -sampling operation returns a tuple chosen from $\sigma_{\mathbf{z}}(\text{Join}(\mathcal{Q}))$ uniformly at random, or returns `nil` if $\sigma_{\mathbf{z}}(\text{Join}(\mathcal{Q})) = \emptyset$.

► **Problem 1.** Preprocess $\mathcal{Q}$ into a data structure that can support $(\mathcal{G}, \mathcal{Z})$ -sampling operations, with the requirement that the sample returned by each operation must be independent of the outputs of all previous operations.

We call the above the $(\mathcal{G}, \mathcal{Z})$ -sampling problem. A data structure is *feasible* if it uses $\tilde{O}(\text{IN})$ space and supports a sampling operation in $\tilde{O}(1)$ time where the notation $\tilde{O}(\cdot)$ hides a factor polylogarithmic to IN .

($\mathcal{G}, \bigvee_{i=1}^m \mathcal{Z}_i$)-Sampling. Fix a schema graph $\mathcal{G} = (\mathcal{V}, \mathcal{E})$ whose size is assumed to be a constant, as well as m subsets $\mathcal{Z}_1, \mathcal{Z}_2, \dots, \mathcal{Z}_m$ of $\mathcal{V}$ for some $m \geq 1$. Let $\mathcal{Q}$ be a join instance of $\mathcal{G}$. Given m tuples $\mathbf{z}_1, \mathbf{z}_2, \dots, \mathbf{z}_m$ where $\mathbf{z}_i$ is over $\mathcal{Z}_i$ for each $i \in [m]$, define

$$\sigma_{\mathbf{z}_1 \vee \dots \vee \mathbf{z}_m}(\text{Join}(\mathcal{Q})) = \bigcup_{i=1}^m \sigma_{\mathbf{z}_i}(\text{Join}(\mathcal{Q})) \quad (3)$$

where $\sigma_{\mathbf{z}_i}(\text{Join}(\mathcal{Q}))$ is as defined in (2). In other words, $\sigma_{\mathbf{z}_1 \vee \dots \vee \mathbf{z}_m}(\text{Join}(\mathcal{Q}))$ includes all and only the tuples $\mathbf{u}$ in the join result that satisfy a condition given in a disjunctive normal form: $\bigvee_{i=1}^m (\bigwedge_{X \in \mathcal{Z}_i} \mathbf{u}(X) = \mathbf{z}_i(X))$. The value of m is permitted to be arbitrarily greater than $2^{|\mathcal{V}|}$, implying that $\mathcal{Z}_1, \mathcal{Z}_2, \dots, \mathcal{Z}_m$ need not be distinct¹.

Given m tuples $\mathbf{z}_1, \mathbf{z}_2, \dots, \mathbf{z}_m$, a $(\mathcal{G}, \bigvee_{i=1}^m \mathcal{Z}_i)$ -sampling operation returns a tuple chosen from $\sigma_{\mathbf{z}_1 \vee \dots \vee \mathbf{z}_m}(\text{Join}(\mathcal{Q}))$ uniformly at random, or returns `nil` if $\sigma_{\mathbf{z}_1 \vee \dots \vee \mathbf{z}_m}(\text{Join}(\mathcal{Q})) = \emptyset$.

► **Problem 2.** Preprocess $\mathcal{Q}$ into a data structure that can support $(\mathcal{G}, \bigvee_{i=1}^m \mathcal{Z}_i)$ -sampling operations, with the requirement that the sample returned by each operation must be independent of the outputs of all previous operations.

We call the above the $(\mathcal{G}, \bigvee_{i=1}^m \mathcal{Z}_i)$ -sampling problem.

1.2 Our Results

Problem 1 (($\mathcal{G}, \mathcal{Z}$)-Sampling). The following is a conjecture that has been extensively used in studying fine-grained complexities; see, e.g., [14, 18, 23, 24, 29, 30] and their references.

Strong Set Disjointness (SSD) Conjecture. Let $S_1, S_2, \dots, S_m$ be $m \geq 2$ sets whose elements are integers. Define $N = \sum_{i=1}^m |S_i|$. In the *set disjointness problem*, we want to preprocess these sets into a data structure of $\mathcal{S}$ space such that, given any distinct integers $a, b \in [m]$, we can report in $\mathcal{T} = o(N)$ time whether $S_a \cap S_b = \emptyset$. The SSD conjecture states that $\mathcal{S}$ must be $\Omega(N^2 / (\mathcal{T}^2 \text{polylog } N))$.

We investigate under the above conjecture the existence of feasible structures for Problem 1. It turns out that the answer depends on whether the schema graph $\mathcal{G} = (\mathcal{V}, \mathcal{E})$ of $\mathcal{Q}$ is “ext- $\mathcal{Z}$ -connex”, which is a notion introduced by Bagan, Durand and Grandjean [5]:

¹ For example: `SELECT * FROM Payment P, TaxPayer T WHERE P.ssn = T.ssn AND (T.job = “prof” OR T.job = “lawyer”)`. Here, both $\mathcal{Z}_1$ and $\mathcal{Z}_2$ are $\{\text{job}\}$.

► **Definition 2.** Let $\mathcal{G} = (\mathcal{V}, \mathcal{E})$ be a schema graph and $\mathcal{Z}$ be a subset of $\mathcal{V}$. We say that $\mathcal{G}$ is **ext- $\mathcal{Z}$ -connex** if the hypergraph $(\mathcal{V}, \mathcal{E} \cup \{\mathcal{Z}\})$ is acyclic.

► **Example 3.** Consider again the schema graph $\mathcal{G} = (\mathcal{V}, \mathcal{E})$ in Example 1. The table below lists the answers to “is $\mathcal{G}$ ext- $\mathcal{Z}$ -connex?” for several representative choices of $\mathcal{Z}$.

$\mathcal{Z}$	$\emptyset$	E	AE	ACE	ACEH	ACEHI	ACEHF
ext- $\mathcal{Z}$ -connex?	yes	yes	no	yes	yes	no	yes

Figure 1(b) shows a join tree of the hypergraph $(\mathcal{V}, \mathcal{E} \cup \{\text{ACEHF}\})$, which serves as evidence that $(\mathcal{V}, \mathcal{E} \cup \{\text{ACEHF}\})$ is acyclic; hence, $\mathcal{G}$ is ext-ACEHF-connex. $\lrcorner$

Our dichotomy result is:

► **Theorem 4.** Let $\mathcal{G} = (\mathcal{V}, \mathcal{E})$ be an acyclic schema graph and $\mathcal{Z}$ be a subset of $\mathcal{V}$. Subject to the SSD-conjecture, the $(\mathcal{G}, \mathcal{Z})$ -sampling problem admits a feasible structure if and only if $\mathcal{G}$ is ext- $\mathcal{Z}$ -connex.

Prior to our work, ext- $\mathcal{Z}$ -connexity has been used to prove dichotomies in several settings [5, 9–13, 21], all of which seem quite different from $(\mathcal{G}, \mathcal{Z})$ -sampling. Interestingly, the dichotomy in Theorem 4 implies an approach to implement feasible $(\mathcal{G}, \mathcal{Z})$ -sampling via “direct-access” queries. To explain, let L be an ordering of the attributes in $\mathcal{V}$: $X_1, X_2, \dots, X_{|\mathcal{V}|}$. This L defines a lexicographic order $\prec$ on the tuples in $\text{Join}(\mathcal{Q})$: $\mathbf{u}_1 \prec \mathbf{u}_2$ if there is an $i \in [|\mathcal{V}|]$ such that $\mathbf{u}_1(X_i) < \mathbf{u}_2(X_i)$ but $\mathbf{u}_1(X_j) = \mathbf{u}_2(X_j)$ for all $1 \leq j \leq i-1$. Given an integer t , a *direct access* (DA) query [6, 7, 12, 13] returns the t -th tuple of $\text{Join}(\mathcal{Q})$ under $\prec$ if $t \in [|\text{Join}(\mathcal{Q})|]$, or nil otherwise. As shown in [12], if $\mathcal{G}$ is ext- $\mathcal{Z}$ -connex, we can find an order L such that

- the attributes of $\mathcal{Z}$ form a prefix of L , i.e., $X_i \in \mathcal{Z}$ for every $i \in [|\mathcal{Z}|]$;
- one can build a structure of $O(\text{IN})$ space that answers any DA query in $O(\log \text{IN})$ time. Fix an arbitrary tuple $\mathbf{z}$ over $\mathcal{Z}$. Crucially, the tuples in $\sigma_{\mathbf{z}}(\text{Join}(\mathcal{Q}))$ must be *consecutive* under $\prec$ (because $\mathcal{Z}$ is a prefix of L). Thus, if $\sigma_{\mathbf{z}}(\text{Join}(\mathcal{Q})) \neq \emptyset$, there exist integers t_1 and t_2 such that a tuple is in $\sigma_{\mathbf{z}}(\text{Join}(\mathcal{Q}))$ if and only if it is the t -th tuple under $\prec$ for some $t \in [t_1, t_2]$. With $O(\log \text{IN})$ DA queries, one can obtain the values of t_1 and t_2 (or certify their non-existence) via binary search. Then, a sample of $\sigma_{\mathbf{z}}(\text{Join}(\mathcal{Q}))$ can be drawn by generating a random $t \in [t_1, t_2]$ and issuing one more DA query. This supports a $(\mathcal{G}, \mathcal{Z})$ -sampling operation in $O(\log^2 \text{IN})$ time.²

In this work, we show that the sample time can be reduced to constant:

► **Theorem 5.** Consider the $(\mathcal{G}, \mathcal{Z})$ -sampling problem where $\mathcal{G}$ is ext- $\mathcal{Z}$ -connex. Given a join instance $\mathcal{Q}$ of $\mathcal{G}$, we can build a structure of $O(\text{IN})$ space in $O(\text{IN})$ expected time such that a $(\mathcal{G}, \mathcal{Z})$ -sampling operation on $\mathcal{Q}$ can be supported in $O(1)$ time.

Problem 2 (($\mathcal{G}, \bigvee_{i=1}^m \mathcal{Z}_i$)-Sampling). Our core contribution to this problem is a new algorithm for a fundamental sampling problem:

Set Union Sampling: Let $S_1, S_2, \dots, S_m$ be $m \geq 2$ sets of elements drawn from a certain domain. Each S_i ($i \in [m]$) supports three operations in constant time:

1. (*size*) return $|S_i|$;
2. (*membership*) check whether a given element is in S_i ;
3. (*sampling*) return an element of S_i chosen uniformly at random.

The goal is to sample an element from $\bigcup_{i=1}^m S_i$ uniformly at random using only these operations. The sample obtained each time must be independent of all previous samples.

² We thank an anonymous reviewer for pointing out the connection to DA queries.

We will prove:

► **Theorem 6.** *There is a set union sampling algorithm with $O(m)$ expected sample time.*

Currently the fastest expected sample time [4, 13] – as will be reviewed in Section 2 – is $O(\min\{m^2, m \log^2 N\})$ where $N = \sum_{i=1}^m |S_i|$, which we strictly improve. By combining Theorems 5 and 6 with additional ideas, we will prove:

► **Theorem 7.** *Consider $(\mathcal{G}, \bigvee_{i=1}^m \mathcal{Z}_i)$ -sampling where $\mathcal{G}$ is ext- $\mathcal{Z}_i$ -connex for all $i \in [m]$. Given a join instance $\mathcal{Q}$ of $\mathcal{G}$, we can build a structure of $O(\text{IN})$ space in $O(\text{IN})$ expected time such that a $(\mathcal{G}, \bigvee_{i=1}^m \mathcal{Z}_i)$ -sampling operation on $\mathcal{Q}$ can be supported in $O(m)$ expected time.*

The connexity requirement is necessary for $m \leq \text{IN}^{0.49}$, as formally stated below:

► **Theorem 8.** *Consider $(\mathcal{G}, \bigvee_{i=1}^m \mathcal{Z}_i)$ -sampling where $m \leq \text{IN}^{0.49}$. If $\mathcal{G}$ is not ext- $\mathcal{Z}_i$ -connex for an arbitrary $i \in [m]$, no structures of $\tilde{O}(\text{IN})$ space can guarantee $\tilde{O}(m)$ expected sample time, subject to the SSD conjecture.*

Random Enumeration. *Random enumeration* of a query result produces a (uniformly) random permutation of the elements therein. An algorithm achieves a *delay* of Δ if it can (i) output the first element or declare an empty result within Δ time, and (ii) after the previous output, produce the next element or declare “no more” within an additional Δ time. The algorithm ensures an *expected* delay Δ if it satisfies the preceding requirement, except that the two “ Δ time” occurrences are replaced with “ Δ expected time”. By combining our sampling techniques with a new enumeration-to-sampling reduction, we prove:

► **Theorem 9.** *Let $\mathcal{G} = (\mathcal{V}, \mathcal{E})$ be a schema graph and $\mathcal{Z}$ be a subset of $\mathcal{V}$ such that $\mathcal{G}$ is ext- $\mathcal{Z}$ -connex. Given a join instance $\mathcal{Q}$ of $\mathcal{G}$, we can build a structure of $O(\text{IN})$ space in $O(\text{IN})$ expected time such that, given any tuple $\mathbf{z}$ over $\mathcal{Z}$, we can randomly enumerate $\sigma_{\mathbf{z}}(\text{Join}(\mathcal{Q}))$ with an expected delay of $O(1)$.*

2 Related Work

At a high level, the objective of join sampling is to create a data structure on the input relations of a join $\mathcal{Q}$ that can be used to extract a uniformly random tuple from $\text{Join}(\mathcal{Q})$. The samples obtained from repetitive extractions must be mutually independent. If $\mathcal{Q}$ is acyclic, Zhao et al. [35] described an $O(\text{IN})$ -space structure ensuring $O(1)$ sample time. The problem becomes more challenging when $\mathcal{Q}$ is cyclic. Improving over [16, 19], Kim et al. [28] presented a structure of $O(\text{IN})$ space that can guarantee a sample time of $O(\text{AGM}/\max\{1, \text{OUT}\})$, where AGM is the join’s *AGM bound* [3], and $\text{OUT} = |\text{Join}(\mathcal{Q})|$. For a broad class of joins with “degree constraints”, Wang and Tao [32] managed to reduce the sample time to $O(\text{polymat}/\max\{1, \text{OUT}\})$, where *polymat* is the *polymatroid bound* [27] of $\mathcal{Q}$, which never exceeds but can be significantly lower than the AGM bound. Recently, Capelli et al. [8] presented a different approach to match the result of [32] up to polylogarithmic factors.

We are not aware of previous work on join sampling under conjunctive or disjunctive selection predicates. Like [35], our study concentrates on acyclic joins, but the sampling problem in [35] is merely a case of $(\mathcal{G}, \mathcal{Z})$ -sampling (see Problem 1) where $\mathcal{Z} = \emptyset$; the result of [35] can be regarded as a special version of our Theorem 5.

The set union sampling problem in Section 1.2 has been studied in two independent articles [4, 13]. Both articles described an algorithm ensuring an expected sample time of $O(m \cdot (N/n))$ where $n = |\bigcup_{i=1}^m S_i|$ and $N = \sum_{i=1}^m |S_i|$ (Section 5.1 will discuss this algorithm

in detail). Note that N can reach mn (this happens when $S_1 = S_2 = \dots = S_m$), in which case the sample time becomes $O(m^2)$. In [4], Aumuller et al. presented another algorithm to extract, in $O(m \log^2 N)$ expected time, a sample that is uniform “with high probability”. Our Theorem 6 strictly improves these results.

Given an acyclic join $\mathcal{Q}$, Carmeli et al. [13] showed how to randomly enumerate $\text{Join}(\mathcal{Q})$ with a delay of $O(\log \text{IN})$, after an $O(\text{IN} \log \text{IN})$ -time preprocessing. In Theorem 9, we settle a more general problem (i.e., random enumeration of $\sigma_z(\text{Join}(\mathcal{Q}))$) with an expected delay of $O(1)$. Given m acyclic joins $\mathcal{Q}_1, \mathcal{Q}_2, \dots, \mathcal{Q}_m$ whose results have the same schema, Carmeli et al. [13] described a way to randomly enumerate $\bigcup_{i=1}^m \text{Join}(\mathcal{Q}_i)$ with an expected delay of $O(m^2 \log \text{IN}_{\text{all}})$ after an $O(\text{IN}_{\text{all}} \log \text{IN}_{\text{all}})$ -time preprocessing, where IN_{all} is the sum of the input sizes of $\mathcal{Q}_1, \dots, \mathcal{Q}_m$. Our set union sampling algorithm in Section 5.1 can be used to reduce their expected delay to $O(m \log \text{IN}_{\text{all}})$. Random enumeration of $\bigcup_{i=1}^m \text{Join}(\mathcal{Q}_i)$ for general (cyclic) joins $\mathcal{Q}_1, \dots, \mathcal{Q}_m$ was considered in [19]; we do not delve into their results further because they are subsumed by those of [13] on acyclic joins (i.e., this work’s focus).

Join reporting (rather than sampling) under selection has been investigated in [18, 34]. To explain their findings, fix a schema graph $\mathcal{G} = (\mathcal{V}, \mathcal{E})$ of a constant size (note: $\mathcal{G}$ can be cyclic), together with a non-empty subset $\mathcal{Z} \subseteq \mathcal{V}$. Consider the task of creating a data structure on a join instance $\mathcal{Q}$ of $\mathcal{G}$ such that, when supplied with a tuple $\mathbf{z}$ over $\mathcal{Z}$, the structure can report $\sigma_z(\text{Join}(\mathcal{Q}))$ efficiently. The focus of [18, 34] is to study the relationships between the space consumption – denoted as $\mathcal{S}$ – of the structure and the time of computing $\sigma_z(\text{Join}(\mathcal{Q}))$. Targeting “output-sensitive” algorithms that report $\sigma_z(\text{Join}(\mathcal{Q}))$ in $\mathcal{T} + O(|\text{OUT}|)$ time where $\text{OUT} = |\sigma_z(\text{Join}(\mathcal{Q}))|$, Zhao et al. [34] derived smooth tradeoffs between $\mathcal{S}$ and $\mathcal{T}$. In Section 6, we complement their results by identifying scenarios where $\mathcal{S} = O(\text{IN})$ and $\mathcal{T} = O(1)$ are possible, provided that the time of reporting $\sigma_z(\text{Join}(\mathcal{Q}))$ is allowed to be $\mathcal{T} + O(|\text{OUT}|)$ *expected*. The work of [18], on the other hand, concerns a different form of tradeoffs that do not bear direct relevance to our results.

3 Preliminaries

3.1 Weighted Sampling

Let S be a set of n elements, denoted as $e_1, e_2, \dots, e_n$, where each e_i ($i \in [n]$) carries a positive integer *weight* $w(e_i)$. Set $W = \sum_{i=1}^n w(e_i)$. A *weighted sampling* operation returns a random element X such that $\Pr[X = e_i] = w(e_i)/W$ for each $i \in [n]$. The *alias method* [31] supports such an operation in $O(1)$ time, after creating a data structure of $O(n)$ space over S in $O(n)$ time. We will refer to the structure as the *alias structure*.

3.2 Acyclic Join Sampling without Selections

This subsection outlines how to build a structure over an acyclic join $\mathcal{Q}$ that uses $O(\text{IN})$ space and can extract a uniformly random tuple from $\text{Join}(\mathcal{Q})$ in $O(1)$ time.

Take a join tree $\mathcal{T}$ of the schema graph $\mathcal{G} = (\mathcal{V}, \mathcal{E})$ of $\mathcal{Q}$. Root $\mathcal{T}$ at an arbitrary node e^* (recall that each node of $\mathcal{T}$ is a hyperedge in $\mathcal{E}$) – doing so enables us to speak about the “parents” and “children” of the nodes in $\mathcal{T}$. Given a node e in $\mathcal{T}$, we denote by $\mathcal{T}_e$ the subtree of $\mathcal{T}$ induced by all the descendants of e (note: e is a descendant of itself). Define

$$\mathcal{Q}_e = \{R_{e'} \mid \text{node } e' \text{ is in } \mathcal{T}_e\} \quad (4)$$

where $R_{e'}$ is the (only) relation in $\mathcal{Q}$ whose schema is e' . We will refer to $\mathcal{Q}_e$ as a *subjoin*. The schema graph of $\mathcal{Q}_e$ is $\mathcal{G}_e = (\mathcal{V}_e, \mathcal{E}_e)$ where $\mathcal{E}_e = \{e' \in \mathcal{E} \mid \text{node } e' \text{ is in } \mathcal{T}_e\}$ and $\mathcal{V}_e = \bigcup_{e' \in \mathcal{E}_e} e'$.

Given a tuple $\mathbf{u} \in R_e$, let us examine $\mathbf{u} \bowtie \text{Join}(\mathcal{Q}_e)$, namely, the result tuples of the subjoin $\mathcal{Q}_e$ to which $\mathbf{u}$ “contributes”. Define the *subjoin weight* of $\mathbf{u}$ as

$$\text{subj-}w_e(\mathbf{u}) = \left| \mathbf{u} \bowtie \text{Join}(\mathcal{Q}_e) \right| \quad (5)$$

where $\mathcal{Q}_e$ is given in (4). The following is implicit from the arguments in [13, 35]. In Appendix A, we present an explicit proof for completeness.

► **Lemma 10.** *For any node e of $\mathcal{T}$, the subjoin weights of all the tuples in R_e can be computed in $O(\text{IN})$ expected time.*

The next lemma, again proved in Appendix A, can be deployed to build a sampling structure at every node of $\mathcal{T}$.

► **Lemma 11.** *Let $\mathcal{Q}$ be an acyclic join with schema graph $\mathcal{G}$, $\mathcal{T}$ be a rooted join tree of $\mathcal{G}$, and e be a node of $\mathcal{T}$. We can build in $O(\text{IN})$ expected time a structure of $O(\text{IN})$ space that, given any $\mathbf{u} \in R_e$, can draw a uniformly random tuple from $\mathbf{u} \bowtie \text{Join}(\mathcal{Q}_e)$ in $O(1)$ time.*

It is now a simple exercise to design the overall structure for sampling from $\text{Join}(\mathcal{Q}) = \text{Join}(\mathcal{Q}_{e^*})$. Notice that $\text{Join}(\mathcal{Q})$ is decomposed into a collection of sets: $\{\mathbf{u} \bowtie \text{Join}(\mathcal{Q}_e) \mid \mathbf{u} \in R_{e^*}\}$. Random sampling from $\text{Join}(\mathcal{Q})$ can be performed in two steps. First, obtain a random tuple $\mathbf{X} \in R_{e^*}$ such that $\Pr[\mathbf{X} = \mathbf{u}] = |\mathbf{u} \bowtie \text{Join}(\mathcal{Q}_{e^*})| / |\text{Join}(\mathcal{Q})| = \text{subj-}w_{e^*}(\mathbf{u}) / |\text{Join}(\mathcal{Q})|$ for each $\mathbf{u} \in R_{e^*}$. As $|\text{Join}(\mathcal{Q})| = \sum_{\mathbf{u} \in R_{e^*}} |\mathbf{u} \bowtie \text{Join}(\mathcal{Q}_{e^*})|$, this is an instance of weighted sampling once the $\text{subj-}w_{e^*}(\mathbf{u})$ value of every $\mathbf{u} \in R_{e^*}$ has been calculated from Lemma 10. An alias structure on R_{e^*} allows us to obtain $\mathbf{X}$ in $O(1)$ time. Second, apply Lemma 11 to draw a uniformly random tuple from $\mathbf{X} \bowtie \text{Join}(\mathcal{Q}_{e^*})$ in $O(1)$ time.

3.3 Properties of Ext- $\mathcal{Z}$ -Connexity

In this subsection, we first introduce several concepts related to ext- $\mathcal{Z}$ -connexity that will aid our exposition in later sections. Then, we review two key properties of ext- $\mathcal{Z}$ -connexity essential for our technical development.

► **Definition 12.** *Let $\mathcal{G} = (\mathcal{V}, \mathcal{E})$ be an acyclic schema graph, $\mathcal{Z}$ be a subset of $\mathcal{V}$, and $\mathcal{T}$ be a join tree of $\mathcal{G}$. We say that $\mathcal{G}$ is **$\mathcal{Z}$ -canonical** if there is a hyperedge $e \in \mathcal{E}$ satisfying $\mathcal{Z} \subseteq e$. An edge $\{e_1, e_2\}$ of $\mathcal{T}$ is **$\mathcal{Z}$ -breakable** if $e_1 \cap e_2 \subseteq \mathcal{Z}$ (note that e_1 and e_2 are hyperedges of $\mathcal{E}$; and $\{e_1, e_2\}$ represents an undirected edge in $\mathcal{T}$).*

► **Definition 13.** *Let $\mathcal{G} = (\mathcal{V}, \mathcal{E})$ be an acyclic schema graph, $\mathcal{Z}$ be a subset of $\mathcal{V}$, and $\mathcal{T}$ be a join tree of $\mathcal{G}$. Suppose that removing all the $\mathcal{Z}$ -breakable edges of $\mathcal{T}$ partitions $\mathcal{T}$ into s subtrees $\mathcal{T}_1, \mathcal{T}_2, \dots, \mathcal{T}_s$. For each $i \in [s]$, define a hypergraph $\mathcal{G}_i = (\mathcal{V}_i, \mathcal{E}_i)$ with $\mathcal{E}_i = \{e \in \mathcal{E} \mid \text{node } e \text{ exists in } \mathcal{T}_i\}$ and $\mathcal{V}_i = \bigcup_{e \in \mathcal{E}_i} e$. Each $\mathcal{G}_i$ is a **$(\mathcal{Z}, \mathcal{T})$ -component** of $\mathcal{G}$.*

Given a schema graph $\mathcal{G} = (\mathcal{V}, \mathcal{E})$, we call two distinct vertices $X, Y \in \mathcal{V}$ *neighbors* if they appear together in at least one hyperedge. A *path* in $\mathcal{G}$ is a sequence of distinct vertices $X_1, X_2, \dots, X_t$ such that X_i and X_{i+1} are neighbors for all $i \in [t-1]$.

► **Definition 14.** *Let $\mathcal{G} = (\mathcal{V}, \mathcal{E})$ be an acyclic schema graph and $\mathcal{Z}$ be a subset of $\mathcal{V}$. A **$\mathcal{Z}$ -path** is a path $Z_1, X_1, X_2, \dots, X_\ell, Z_2$ (where $\ell \geq 1$) such that (i) Z_1 and Z_2 belong to $\mathcal{Z}$, but they are not neighbors; (ii) $X_i \notin \mathcal{Z}$ for every $i \in [\ell]$.*

► **Example 15.** Consider the schema graph $\mathcal{G} = (\mathcal{V}, \mathcal{E})$ in Figure 1(a). Among the choices of $\mathcal{Z}$ in Example 3, $\mathcal{G}$ is $\mathcal{Z}$ -canonical only for $\mathcal{Z} = \emptyset$ and $\mathcal{Z} = \mathbf{E}$.

Set $\mathcal{Z}$ to ACEHF. The edge $\{ABC, CDE\}$ of $\mathcal{T}$ is $\mathcal{Z}$ -breakable because $ABC \cap CDE = C$ is a subset of $\mathcal{Z}$. On the other hand, the edge $\{EGH, GI\}$ is not $\mathcal{Z}$ -breakable because $EGH \cap GI = G$ is not a subset of $\mathcal{Z}$. By removing the $\mathcal{Z}$ -breakable edges, we obtain four subtrees of $\mathcal{T}$: the first contains nodes e_1 and e_3 , while every other node (i.e., e_2 , e_4 , and e_5) forms a subtree by itself. Accordingly, $\mathcal{G}$ has four $(\mathcal{Z}, \mathcal{T})$ -components: $\mathcal{G}_1 = (EGHI, \{EGH, GI\})$, $\mathcal{G}_2 = (ABC, \{ABC\})$, $\mathcal{G}_3 = (CDE, \{CDE\})$, and $\mathcal{G}_4 = (EFJ, \{EFJ\})$. There are no $\mathcal{Z}$ -paths in $\mathcal{G}$.

Set $\mathcal{Z}$ instead to ACEHI. The path A, B, D, G, I is a $\mathcal{Z}$ -path. $\lrcorner$

► **Lemma 16** ([5, Lemma 23]). *Let $\mathcal{G} = (\mathcal{V}, \mathcal{E})$ be an acyclic schema graph and $\mathcal{Z}$ be a subset of $\mathcal{V}$. The following statements are equivalent:*

1. $\mathcal{G}$ is ext- $\mathcal{Z}$ -connex.
2. $\mathcal{G}$ does not have a $\mathcal{Z}$ -path.
3. Let $\mathcal{T}$ be an arbitrary join tree of $\mathcal{G}$. Denote by s the number of $(\mathcal{Z}, \mathcal{T})$ -components of $\mathcal{G}$, and by $\mathcal{G}_i = (\mathcal{V}_i, \mathcal{E}_i)$ the i -th component. For every $i \in [s]$, $\mathcal{G}_i$ is $(\mathcal{V}_i \cap \mathcal{Z})$ -canonical.

► **Example 17.** Continuing Example 15, for $\mathcal{Z} = \text{ACEHF}$, as mentioned $\mathcal{G}$ has $(\mathcal{Z}, \mathcal{T})$ -components: $\mathcal{G}_1 = (ABC, \{ABC\})$, $\mathcal{G}_2 = (CDE, \{CDE\})$, $\mathcal{G}_3 = (EFJ, \{EFJ\})$, and $\mathcal{G}_4 = (EGHI, \{EGH, GI\})$. For every $1 \leq i \leq 4$, $\mathcal{G}_i$ is $(\mathcal{V}_i \cap \mathcal{Z})$ -canonical; hence, $\mathcal{G}$ is ext- $\mathcal{Z}$ -connex. However, for $\mathcal{Z} = \text{ACEHI}$, as mentioned there is a $\mathcal{Z}$ -path in $\mathcal{G}$, which is thus not ext- $\mathcal{Z}$ -connex. $\lrcorner$

4 A Dichotomy on $(\mathcal{G}, \mathcal{Z})$ -Sampling

Theorem 4 involves a negative result (the “only-if direction”) and a positive result (the “if direction”). We will prove the negative result in Section 4.1. The positive result is a corollary of Theorem 5, whose proof is presented in Section 4.2.

4.1 The Only-If Direction of Theorem 4

As explained in Section 1.2, the input to the set disjointness problem comprises $m \geq 2$ sets $S_1, \dots, S_m$. Given distinct integers $a, b \in [m]$, a *disjointness query* returns whether $S_a \cap S_b = \emptyset$. The SSD-conjecture states that any structure promising to answer such a query in $O(\text{polylog } N)$ time must use $\Omega(N^2 / \text{polylog } N)$ space, where $N = \sum_{i=1}^m |S_i|$.

Suppose that the only-if claim of Theorem 4 is false. Thus, there exist a hypergraph $\mathcal{G} = (\mathcal{V}, \mathcal{E})$ and a subset $\mathcal{Z} \subseteq \mathcal{V}$ such that $\mathcal{G}$ is not ext- $\mathcal{Z}$ -connex but a feasible $(\mathcal{G}, \mathcal{Z})$ -sampling structure exists. We will show how to build a set-disjointness structure of $\tilde{O}(N)$ space that answers a disjointness query in $\tilde{O}(1)$ time, thus breaking the SSD-conjecture. Since $\mathcal{G}$ is not ext- $\mathcal{Z}$ -connex, it must have a $\mathcal{Z}$ -path by Lemma 16. Pick an arbitrary $\mathcal{Z}$ -path: $A, X_1, X_2, \dots, X_\ell, B$ for some $\ell \geq 1$. Remember that no vertex on the path belongs to $\mathcal{Z}$ except A and B .

From $S_1, \dots, S_m$ (the input to set disjointness), next we create a join $\mathcal{Q}$ whose schema graph is $\mathcal{G} = (\mathcal{V}, \mathcal{E})$. For each hyperedge $e \in \mathcal{E}$, construct a relation R_e with schema e as follows. For each integer $i \in [m]$ and every element $x \in S_i$, insert a tuple $\mathbf{u}$ into R_e (provided that $\mathbf{u}$ is not already in R_e) such that (i) if $A \in e$, then $\mathbf{u}(A) = i$; (ii) if $B \in e$, then $\mathbf{u}(B) = i$; (iii) for every attribute $Z \in (e \cap \mathcal{Z}) \setminus \{A, B\}$, set $\mathbf{u}(Z) = \perp$ where $\perp$ is a special symbol; (iv) for every attribute $X \in e \setminus \mathcal{Z}$, set $\mathbf{u}(X) = x$. By Definition 14, no hyperedge of $\mathcal{E}$ covers both A and B ; hence, at most one step between (i) and (ii) applies. The relation R_e thus designed has a size at most N .

► **Lemma 18.** *Given distinct $a, b \in [m]$, define a tuple $\mathbf{z}$ over $\mathcal{Z}$ with $\mathbf{z}(A) = a$, $\mathbf{z}(B) = b$, and $\mathbf{z}(Z) = \perp$ for every $Z \in \mathcal{Z} \setminus \{A, B\}$. Then, $S_a \cap S_b \neq \emptyset$ if and only if $\sigma_{\mathbf{z}}(\text{Join}(\mathcal{Q})) \neq \emptyset$.*

Proof. Let us first prove the “if direction” ($\Leftarrow$). Consider an arbitrary tuple $\mathbf{u} \in \sigma_{\mathbf{z}}(\text{Join}(\mathcal{Q}))$. Clearly, $\mathbf{u}(A) = a$ and $\mathbf{u}(B) = b$. Since $A, X_1, X_2, \dots, X_\ell, B$ is a $\mathcal{Z}$ -path, there exist hyperedges $e_0, e_1, \dots, e_\ell$ in $\mathcal{E}$ such that (i) $\{A, X_1\} \subseteq e_0$, (ii) $\{X_i, X_{i+1}\} \subseteq e_i$ for each $i \in [\ell - 1]$, and (iii) $\{X_\ell, B\} \subseteq e_\ell$. We argue:

$$\mathbf{u}(X_1) = \mathbf{u}(X_2) = \dots = \mathbf{u}(X_\ell). \quad (6)$$

To see why, fix any $i \in [\ell - 1]$. As $\mathbf{u} \in \sigma_{\mathbf{z}}(\text{Join}(\mathcal{Q}))$, we know $\mathbf{u}[e_i] \in R_{e_i}$. As X_i and X_{i+1} are attributes outside Z , our construction of $\mathcal{Q}$ ensures $\mathbf{u}(X_i) = \mathbf{u}(X_{i+1})$. The fact that this holds for every $i \in [\ell - 1]$ proves (6). By the construction of $\mathcal{Q}$, $\mathbf{u}(X_1)$ is an element of S_a and $\mathbf{u}(X_\ell)$ is an element of S_b . Therefore, (6) tells us $S_a \cap S_b \neq \emptyset$.

Next, we prove the “only if direction” ($\Rightarrow$). Fix any $x \in S_a \cap S_b$. Construct a tuple $\mathbf{u}$ over $\mathcal{V}$ where $\mathbf{u}(A) = a, \mathbf{u}(B) = b, \mathbf{u}(Z) = \perp$ for every $Z \in \mathcal{Z} \setminus \{A, B\}$, and $\mathbf{u}(X) = x$ for every $X \in \mathcal{V} \setminus \mathcal{Z}$. It suffices to prove that $\mathbf{u} \in \sigma_{\mathbf{z}} \text{Join}(\mathcal{Q})$. Indeed, this is true because, for every $e \in \mathcal{E}$, our construction explicitly inserts tuple $\mathbf{u}[e]$ into R_e , noticing that e cannot contain both A and B . $\blacktriangleleft$

The above lemma can be used to break the SSD-conjecture. Note that the input size of $\mathcal{Q}$ satisfies $\text{IN} \leq N \cdot |\mathcal{E}| = O(N)$, thus allowing us to build a structure of $\tilde{O}(\text{IN}) = \tilde{O}(N)$ space on $\mathcal{Q}$ to support a $(\mathcal{G}, \mathcal{Z})$ -sampling operation in $\tilde{O}(1)$ time. A disjointness query parameterized by integers $a, b \in [m]$ can be answered as follows. First, create a tuple $\mathbf{z}$ over $\mathcal{Z}$ such that $\mathbf{z}(A) = a, \mathbf{z}(B) = b$, and $\mathbf{z}(Z) = \perp$ for every $Z \in \mathcal{Z} \setminus \{A, B\}$. Then, issue a $(\mathcal{G}, \mathcal{Z})$ -sampling operation and declare $S_a \cap S_b = \emptyset$ if and only if the operation returns nothing. This correctly answers the query in $\tilde{O}(1)$ time.

We note that similar constructions were used in [5, 7] to prove dichotomies for settings different from ours. Our contribution lies in establishing connections to $(\mathcal{G}, \mathcal{Z})$ -sampling.

4.2 Proof of Theorem 5 (a.k.a. The If-Direction of Theorem 4)

Let $\mathcal{G} = (\mathcal{V}, \mathcal{E})$ be a schema graph and $\mathcal{Z}$ be a subset of $\mathcal{V}$. Given a join instance $\mathcal{Q}$ of $\mathcal{G}$, we want to build a structure of $O(\text{IN})$ space in $O(\text{IN})$ expected time to support a $(\mathcal{G}, \mathcal{Z})$ -sampling operation on $\mathcal{Q}$ in constant time.

Case 1: $\mathcal{G}$ Is $\mathcal{Z}$ -Canonical. By Definition 12, in this case there is a hyperedge $e^* \in \mathcal{E}$ satisfying $\mathcal{Z} \subseteq e^*$. Now, root $\mathcal{T}$ at e^* , apply Lemma 10 to calculate the subjoin weights of all tuples in R_{e^*} (the relation in $\mathcal{Q}$ with schema e^*), and build a sampling structure of Lemma 11 by setting the symbol e there to e^* . Given a tuple $\mathbf{z}$ over $\mathcal{Z}$, define

$$R_{e^*}(\mathbf{z}) = \{\mathbf{u} \in R_{e^*} \mid \mathbf{u}[\mathcal{Z}] = \mathbf{z}[\mathcal{Z}]\}. \quad (7)$$

It is easy to verify

$$\sigma_{\mathbf{z}}(\text{Join}(\mathcal{Q})) = \bigcup_{\mathbf{u} \in R_{e^*}(\mathbf{z})} \mathbf{u} \bowtie \text{Join}(\mathcal{Q}). \quad (8)$$

Random sampling from $\sigma_{\mathbf{z}}(\text{Join}(\mathcal{Q}))$ can be done in a two-step approach similar to what was described at the end of Section 3.2. First, obtain a random tuple $\mathbf{X} \in R_{e^*}(\mathbf{z})$ such that, for each $\mathbf{u} \in R_{e^*}(\mathbf{z})$, we have

$$\Pr[\mathbf{X} = \mathbf{u}] = |\mathbf{u} \bowtie \text{Join}(\mathcal{Q}_{e^*})| / |\sigma_{\mathbf{z}}(\text{Join}(\mathcal{Q}))| = \text{subj-}w_{e^*}(\mathbf{u}) / |\sigma_{\mathbf{z}}(\text{Join}(\mathcal{Q}))|$$

where $\text{subj-}w_{e^*}(\mathbf{u})$ is the subjoin weight of $\mathbf{u}$ (with respect to the rooted $\mathcal{T}$); see (5). As $|\sigma_{\mathbf{z}}(\text{Join}(\mathcal{Q}))| = \sum_{\mathbf{u} \in R_{e^*}(\mathbf{z})} |\mathbf{u} \bowtie \text{Join}(\mathcal{Q}_{e^*})| = \sum_{\mathbf{u} \in R_{e^*}(\mathbf{z})} \text{subj-}w_{e^*}(\mathbf{u})$, this is an instance of weighted sampling since the $\text{subj-}w_{e^*}(\mathbf{u})$ value of every $\mathbf{u} \in R_{e^*}(\mathbf{z})$ is already available.

Hence, an alias structure on $R_{e^*}(z)$ allows us to obtain $\mathbf{X}$ in constant time. Second, apply Lemma 11 to draw a uniformly random tuple from $\mathbf{X} \bowtie \text{Join}(\mathcal{Q}_{e^*})$ in constant time.

The alias structure on $R_{e^*}(z)$ occupies $O(|R_{e^*}(z)|)$ space and can be built in $O(|R_{e^*}(z)|)$ time, as explained in Section 3.1. We do this for every $z \in \Pi_{\mathcal{Z}}(R_{e^*})$, where Π is projection in relational algebra. For distinct $z, z' \in \Pi_{\mathcal{Z}}(R_{e^*})$, the sets $R_{e^*}(z)$ and $R_{e^*}(z')$ are disjoint. Hence, in total, all the alias structures occupy $O(|R_{e^*}|) = O(\text{IN})$ space and can be built in $O(\text{IN})$ expected time (the time is expected because hashing is required to obtain the sets $R_{e^*}(z)$ of each $z \in \Pi_{\mathcal{Z}}(R_{e^*})$).

► **Remark.** We make an observation here that will be useful in Section 5.2. For each $z \in \Pi_{\mathcal{Z}}(R_{e^*})$, as mentioned $|\sigma_z(\text{Join}(\mathcal{Q}))| = \sum_{u \in R_{e^*}(z)} \text{subj-}w_{e^*}(u)$. Since the subjoin weights of all tuples in R_{e^*} are available, we can compute the sizes $|\sigma_z(\text{Join}(\mathcal{Q}))|$ of all $z \in \Pi_{\mathcal{Z}}(R_{e^*})$ by scanning R_{e^*} once in $O(\text{IN})$ time. Storing all those sizes takes $O(|R_{e^*}|) = O(\text{IN})$ extra space. As the benefit, given any tuple z over $\mathcal{Z}$, we can obtain $|\sigma_z(\text{Join}(\mathcal{Q}))|$ in $O(1)$ time with hashing.

Case 2: $\mathcal{G}$ Is Not $\mathcal{Z}$ -Canonical. Let $\mathcal{T}$ be an arbitrary join tree of $\mathcal{G}$. As $\mathcal{G}$ is ext- $\mathcal{Z}$ -connex, by Lemma 16, $\mathcal{T}$ defines a number $s \geq 1$ of $(\mathcal{Z}, \mathcal{T})$ -components of $\mathcal{G}$ – denoted as $\mathcal{G}_1 = (\mathcal{V}_1, \mathcal{E}_1), \dots, \mathcal{G}_s = (\mathcal{V}_s, \mathcal{E}_s)$, respectively – such that $\mathcal{G}_i$ is $(\mathcal{V}_i \cap \mathcal{Z})$ -canonical for every $i \in [s]$. For each $i \in [s]$, define $\mathcal{Z}_i = \mathcal{V}_i \cap \mathcal{Z}$ and $\mathcal{Q}_i = \{R_e \in \mathcal{Q} \mid e \in \mathcal{E}_i\}$. Note that $\mathcal{Q}_i$ is a join instance of the schema graph $\mathcal{G}_i$, which is $\mathcal{Z}_i$ -canonical. Denote by IN_i the input size of $\mathcal{Q}_i$; we have $\sum_{i=1}^s \text{IN}_i = \text{IN}$ because every relation of $\mathcal{Q}$ appears in exactly one of $\mathcal{Q}_1, \dots, \mathcal{Q}_s$. For each $i \in [s]$, we build a $(\mathcal{G}_i, \mathcal{Z}_i)$ -sampling structure on $\mathcal{Q}_i$ in the way explained for Case 1. All the s structures occupy $O(\text{IN})$ space in total and can be built in $O(\text{IN})$ expected time.

Consider now a $(\mathcal{G}, \mathcal{Z})$ -sampling operation, which is given a tuple z over $\mathcal{Z}$. For each $i \in [s]$, defining $z_i = z[\mathcal{V}_i]$, we perform a $(\mathcal{G}_i, \mathcal{Z}_i)$ -sampling operation using z_i ; let us assume that this operation returns u_i . If u_i is **nil** for any i , we return **nil** for the original $(\mathcal{G}, \mathcal{Z})$ -sampling operation. Otherwise, return $u_1 \bowtie u_2 \bowtie \dots \bowtie u_s$, which, as explained in Appendix B, must be a uniformly random tuple of $\sigma_z(\text{Join}(\mathcal{Q}))$. The $(\mathcal{G}, \mathcal{Z})$ -sampling operation takes $O(s) = O(1)$ time overall. This concludes the proof of Theorem 5.

5 $(\mathcal{G}, \bigvee_{i=1}^m \mathcal{Z}_i)$ -Sampling Algorithms

The crux of our solution to Problem 2 is an algorithm optimally settling the set union sampling problem (see Section 1.2). We will present this algorithm in Section 5.1 and explain in Section 5.2 how it leads to a structure for Problem 2 that establishes Theorem 7. Finally, Section 5.3 gives the proof of Theorem 8.

5.1 Set Union Sampling

Recall from Section 1.2 that, in this problem, we have a collection of $m \geq 2$ sets $S_1, S_2, \dots, S_m$, each supporting three $O(1)$ -time operations: *size*, *membership*, and *sample*. The objective is to draw an element from $\bigcup_{i=1}^m S_i$ uniformly at random.

As before, set $n = |\bigcup_{i=1}^m S_i|$ and $N = \sum_{i=1}^m |S_i|$. For each element $y \in \bigcup_{i=1}^m S_i$, define its *inverted set* as $\text{InvS}(y) = \{i \in [m] \mid y \in S_i\}$, i.e., the “ids” of the sets containing y . The *degree* of y is $\deg(y) = |\text{InvS}(y)|$. Let us first describe a baseline method, which was given explicitly in [4] and implicitly in [13], before presenting our new ideas.

Baseline Method. Draw a random value $X \in [m]$ such that $\Pr[X = i] = |S_i|/N$ for each $i \in [m]$. Then, use the sampling operation to draw an element Y from S_X . Finally, carry out an *acceptance step*:

Acceptance Step: Accept Y with probability $1/\deg(Y)$.

If accepted, Y is returned; otherwise, the algorithm repeats from scratch.

The algorithm correctly returns a uniform sample of $\bigcup_{i=1}^m S_i$. To see why, fix an arbitrary element $y \in \bigcup_{i=1}^m S_i$. This element is returned if and only if three conditions are satisfied: (i) $X \in \text{InvS}(y)$, which occurs with probability $|S_X|/N$ for each $X \in \text{InvS}(y)$, (ii) $Y = y$, which occurs with probability $1/|S_X|$ conditioned on X , and (iii) y is accepted, which occurs with probability $1/\deg(y)$ conditioned on y . Hence, the algorithm outputs y with probability

$$\sum_{X \in \text{InvS}(y)} \frac{|S_X|}{N} \cdot \frac{1}{|S_X|} \cdot \frac{1}{\deg(y)} = \frac{1}{N} \quad (9)$$

which is identical for all $y \in \bigcup_{i=1}^m S_i$. As a corollary, in each repeat, the algorithm *succeeds* in returning a sample with probability n/N . Thus, N/n repeats are needed in expectation.

The random value X can be easily obtained in $O(m)$ time. One (logically simple) way to do so is to build an alias structure (see Section 3.2) on $|S_1|, |S_2|, \dots, |S_m|$ – namely, the m set sizes – on the fly in $O(m)$ time (using the size operation), after which X can be extracted from the structure in $O(1)$ time. The time to obtain Y is $O(1)$ (using the sampling operation). The troublemaker, however, is the acceptance step. The standard approach [4, 13] is to query the membership of Y in each S_i ($i \in [m]$), which costs $O(m)$ time. This renders the overall sample time $O(mN/n)$ in expectation.

New Idea: Total Law of Expectation. Our objective is to implement the acceptance step in $O(mn/N)$ expected time – note that this is faster than $O(m)$ by a factor of N/n , which eventually allows us to bring the expected sample time from $O(mN/n)$ down to $O(m)$.

Let us start with a fundamental fact:

► **Proposition 19.** Let Γ be a random variable taking values from $[m]$, and U be a uniformly random variable over $[m]$. If U and Γ are independent, then $\Pr[U \leq \Gamma] = \frac{1}{m} \mathbf{E}[\Gamma]$.

The proposition is the total law of expectation in disguise; see Appendix C for a proof. Equipped with the above, implementing the acceptance step boils down to:

Given $x \in [m]$ and $y \in S_x$, generate a rand. var. $\Gamma \in [m]$ with $\mathbf{E}[\Gamma] = m/\deg(y)$.

Later, we will explain how to achieve the above purpose in $O(m/\deg(y))$ expected time. Once done, we can perform the acceptance step as follows (recall that, prior to this, the baseline method has obtained the values of two random variables X and Y): (i) generate a uniformly random variable U over $[m]$ in constant time; (ii) generate the aforementioned random variable Γ (setting x and y to the values of X and Y , respectively) in $O(m/\deg(Y))$ expected time; (iii) accept if $U \leq \Gamma$.

Correctness follows from Proposition 19 (the acceptance probability is $\frac{1}{m} \mathbf{E}[\Gamma] = 1/\deg(Y)$, as desired). For a particular $x \in [m]$ and a particular $y \in S_x$, we have $\Pr[X = x, Y = y] = \frac{|S_x|}{N} \cdot \frac{1}{|S_x|} = 1/N$. Thus, the expected cost of the acceptance step will be at the order of

$$\begin{aligned} \sum_{x \in [m]} \sum_{y \in S_x} \Pr[X = x, Y = y] \frac{m}{\deg(y)} &= \sum_{x \in [m]} \sum_{y \in S_x} \frac{1}{N} \frac{m}{\deg(y)} \\ &= \frac{m}{N} \sum_{y \in \bigcup_{i=1}^m S_i} \sum_{x \in \text{InvS}(y)} \frac{1}{\deg(y)} = \frac{mn}{N} \end{aligned} \quad (10)$$

where the last step used $\deg(y) = |\text{InvS}(y)|$ and $n = |\bigcup_{i=1}^m S_i|$.

Generation of Γ . Next, we will concentrate on the Γ -generating task defined earlier. Recall that the generation is based on a given set “id” $x \in [m]$ and an element $y \in S_x$. Consider the procedure below:

```

find-2nd ( $x, y$ )
1.  $F \leftarrow 0$  and  $I \leftarrow [m] \setminus \{x\}$ 
2. while  $I \neq \emptyset$  do
3.   sample without replacement (WoR) an integer  $i \in I$ 
   /* note: the WoR-sampling operation removes  $i$  from  $I$  */
4.   if  $y \in S_i$  then return  $F$  else increase  $F$  by 1
5. return  $F$  /* note:  $F$  must be  $m - 1$  here */

```

In plain words, if $\deg(y) \geq 2$, the value F returned is a random variable giving the number of *failed* WoR-sampling operations before finding *another* set $S_i \neq S_x$ containing y ; otherwise, $F = m - 1$. We prove in Appendix C:

► **Proposition 20.** $\mathbf{E}[F] = \frac{m}{\deg(y)} - 1$.

Combining the above with the obvious fact that $F \in [0, m - 1]$, we can now define $\Gamma = F + 1$ as the desired random variable satisfying $\Gamma \in [m]$ and $\mathbf{E}[\Gamma] = m / \deg(y)$. The WoR-sampling operation at Line 3 can be implemented in $O(1)$ time; see, e.g., [13]. Thus, the cost of **find-2nd** is proportional to the value F returned. It follows from Proposition 20 that the expected cost of **find-2nd** is $O(\mathbf{E}[F]) = O(m / \deg(y))$.

Total Cost of Set-Union Sampling. Recall that, in the (original) baseline algorithm, each repeat takes $O(m)$ time; as the number of repeats is N/n expected, the total cost of taking one sample from $\bigcup_{i=1}^m S_i$ is $O(mN/n)$. By applying our remedy, one repeat of the baseline algorithm is now carried out in $O(mn/N)$ expected time (see the analysis in (10)). The remedy does not affect the expected number of repeats (i.e., N/n), which suggests that the overall expected sample time should be $O(\frac{mn}{N} \cdot \frac{N}{n}) = O(m)$. Indeed, if $\mathbf{E}[T_{total}]$ represents the expected time for our algorithm to acquire a sample from $\bigcup_{i=1}^m S_i$, we can write $\mathbf{E}[T_{total}] = O(mn/N) + (1 - n/N) \mathbf{E}[T_{total}]$, which solves to $\mathbf{E}[T_{total}] = O(m)$. This completes the proof of Theorem 6.

► **Remark.** The proposed algorithm is reminiscent of that of Karp, Luby, and Madras [26], which was designed to estimate $|\bigcup_{i=1}^m S_i|$ rather than to sample from $\bigcup_{i=1}^m S_i$. It can, in fact, be incorporated into their framework to perform the estimation in a slightly simpler way. Our key novelty lies in employing Proposition 19 for the acceptance step, while Karp, Luby, and Madras [26] used a more sophisticated method that is tailored for estimation and does not extend to sampling.

5.2 A Structure for Problem 2

Denote by $\mathcal{Q}$ a join instance of a schema graph $\mathcal{G} = (\mathcal{V}, \mathcal{E})$. We prove Theorem 7 on the $(\mathcal{G}, \bigvee_{i=1}^m \mathcal{Z}_i)$ -sampling problem with a reduction to set union sampling. Recall that $\mathcal{G}$ is ext- $\mathcal{Z}_i$ -connex for every $i \in [m]$. Hence, we can create a structure Υ_i of Theorem 5 on every $(\mathcal{G}, \mathcal{Z}_i)$ that supports $(\mathcal{G}, \mathcal{Z}_i)$ -sampling in $O(1)$ time. As each structure occupies $O(\text{IN})$ space, it may appear *as if* the total space would be $O(m \cdot \text{IN})$. This is not true. To see why, note that every $\mathcal{Z}_i$ is a subset of $\mathcal{V}$; as $\mathcal{V}$ has a constant size, the number of its subsets is bounded by a constant. This means that there can be only a constant number of *distinct* subsets among $\mathcal{Z}_1, \mathcal{Z}_2, \dots, \mathcal{Z}_m$. Physically, one structure of Theorem 5 suffices for each distinct subset; hence, the total space is $O(\text{IN})$. For similar reasons, all the structures can be constructed in $O(\text{IN})$ expected time.

A $(\mathcal{G}, \bigvee_{i=1}^m \mathcal{Z}_i)$ -sampling operation is given m tuples $z_1, \dots, z_m$ over $\mathcal{Z}_1, \dots, \mathcal{Z}_m$, respectively. For each $i \in [m]$, define $S_i = \sigma_{z_i}(\text{Join}(\mathcal{Q}))$. The operation, essentially, aims to return a uniformly-random tuple from $\bigcup_{i=1}^m S_i$. To cast this as an instance of set union sampling, we need to implement each of the size, membership, and sampling operations in $O(1)$ time on each S_i . This is trivial for two operations:

- sampling: use Υ_i to perform a $(\mathcal{G}, \mathcal{Z}_i)$ -sampling operation;
- membership: given a tuple u over $\mathcal{V}$, this operation checks whether $u \in \sigma_{z_i}(\text{Join}(\mathcal{Q}))$. This requires checking if $u[e]$ is in the relevant relation of $\mathcal{Q}$ for each $e \in \mathcal{E}$ and if $u(Z) = z_i(Z)$ for each $Z \in \mathcal{Z}_i$. All the checking can be done in $O(1)$ time.

We can support also the size operation in $O(1)$ time by using directly Υ_i . The ideas are similar to those explained in Section 4.2 and deferred to Appendix C. Our algorithm in Theorem 6 can now be utilized to extract a uniformly-random tuple from $\bigcup_{i=1}^m S_i$ in $O(m)$ expected time, thus completing the proof of Theorem 7.

5.3 Proof of Theorem 8

The SSD conjecture given in Section 1.2 concerns data structures that solve the set disjointness problem with *deterministic* query time (here, a “query” is given $a, b \in [m]$ and reports whether $S_a \cap S_b = \emptyset$). To prove Theorem 8, we will first argue that a similar conjecture still stands even on set-disjointness structures with *expected* query time.

Expected SSD Conjecture. Let $S_1, S_2, \dots, S_m$ be $m \geq 2$ sets whose elements are integers. Define $N = \sum_{i=1}^m |S_i|$. We want to preprocess these sets into a data structure of $\mathcal{S}$ space such that, given any distinct integers $a, b \in [m]$, we can report in $\mathcal{T} = o(N)$ expected time whether $S_a \cap S_b = \emptyset$. The *expected SSD conjecture* states that $\mathcal{S}$ must be $\Omega(N^2/(\mathcal{T}^2 \text{polylog } N))$.

► **Lemma 21.** *The SSD conjecture implies the expected SSD conjecture.*

Proof. Suppose that we can build a structure Υ of $\mathcal{S}$ space and $\mathcal{T}$ expected query time for the set disjointness problem. We will prove the existence of a set-disjointness structure of $\mathcal{S} + O(\mathcal{T} \log N)$ space and $O(\mathcal{T} \log N)$ deterministic query time. This establishes the claim in Lemma 21 because $\mathcal{T} \leq N \leq \mathcal{S}$, implying $\mathcal{S} + O(\mathcal{T} \log N) = O(\mathcal{S} \log N)$.

Fix any distinct integers $a, b \in [m]$. Let $\mathcal{A}$ be the algorithm associated with $\mathcal{T}$ for deciding whether $S_a \cap S_b = \emptyset$. Denote by X the cost for $\mathcal{A}$. Note that X is a random variable with $\mathbf{E}[X] \leq \mathcal{T}$. By Markov’s inequality, $\mathbf{Pr}[X \geq 2\mathcal{T}] \leq 1/2$. Imagine running $\mathcal{A}$ for $3 \log N$ times. The probability for all those runs to take at least $2\mathcal{T}$ time to terminate is at most $1/N^3$. In other words, with probability at least $1 - 1/N^3$, at least one run finishes within $2\mathcal{T}$ time.

In general, a randomized algorithm becomes deterministic once all the random bits are fixed. “Running $\mathcal{A}$ ” is equivalent to (i) first fixing a sequence σ of random bits, and (ii) then executing *deterministically* the instructions of $\mathcal{A}$ based on σ . If $\mathcal{A}$ finishes within $2\mathcal{T}$ time, it consumes at most $2\mathcal{T}$ words of random bits. Now, take $t = 3 \log N$ random bit sequences $\sigma_1, \dots, \sigma_t$, each of which is $2\mathcal{T}$ words long. For each $i \in [t]$, define $\mathcal{A}(\sigma_i)$ as the deterministic algorithm that runs $\mathcal{A}$ based on σ_i , with the modification that $\mathcal{A}(\sigma_i)$ terminates itself after having run $\mathcal{A}$ for a duration of $2\mathcal{T}$ time. As per our earlier discussion, with probability at least $1 - 1/N^3$, at least one of $\mathcal{A}(\sigma_1), \dots, \mathcal{A}(\sigma_t)$ manages to report whether $S_a \cap S_b = \emptyset$.

Now, let us consider all $m(m-1)/2$ distinct pairs of $a, b \in [m]$. With probability at least $1 - m^2/N^3 \geq 1 - 1/N$, for every pair of a and b , at least one of $\mathcal{A}(\sigma_1), \dots, \mathcal{A}(\sigma_t)$ manages to report whether $S_a \cap S_b = \emptyset$ – let us call $\{\sigma_1, \dots, \sigma_t\}$ a *working set*. As the probability $1 - 1/N$ is greater than 0 (because $N \geq m \geq 2$), a working set must exist.

Our final set-disjointness structure consists of Υ and a working set of random-bit sequences $\{\sigma_1, \dots, \sigma_t\}$, the total space of which is $\mathcal{S} + O(\mathcal{T} \log N)$. Given two distinct integers $a, b \in [m]$, we run all of $\mathcal{A}(\sigma_1), \dots, \mathcal{A}(\sigma_t)$, at least one of which manages to report whether $S_a \cap S_b = \emptyset$. The total query time is $O(\mathcal{T} \cdot t) = O(\mathcal{T} \log N)$. $\blacktriangleleft$

Our argument in Section 4.1 is a reduction from set disjointness to $(\mathcal{G}, \mathcal{Z})$ -sampling where $\mathcal{G}$ is not ext- $\mathcal{Z}$ -connex, and works regardless of whether the $(\mathcal{G}, \mathcal{Z})$ -sampling algorithm is randomized or not. The reduction shows that if there is a structure of $O(\text{IN})$ space that can support a $(\mathcal{G}, \mathcal{Z})$ -sampling operation in $O(\text{IN}^{0.49})$ expected time, then there is a set-disjointness structure of $O(N)$ space and $O(N^{0.49})$ expected query time – this will break the expected SSD conjecture.

Let us return to the context of Theorem 8 where $m \leq \text{IN}^{0.49}$. W.l.o.g., assume that $\mathcal{G}$ is not ext- $\mathcal{Z}_1$ -connex, and yet there is a structure Υ of $O(\text{IN})$ space that can support a $(\mathcal{G}, \bigvee_{i=1}^m \mathcal{Z}_i)$ -sampling operation in $O(m)$ expected time. We will show how to use Υ to support a $(\mathcal{G}, \mathcal{Z}_1)$ -sampling operation in $O(\text{IN}^{0.49})$ expected time, which, as pointed out earlier, breaks the expected SSD conjecture. In fact, this is fairly obvious. Suppose that we are given a tuple $\mathbf{z}_1$ over $\mathcal{Z}_1$. For each $i \in [2, m]$, construct a dummy tuple $\mathbf{z}_i$ over $\mathcal{Z}_i$ such that, for each $Z \in \mathcal{Z}_i$, $\mathbf{z}_i(Z)$ is a value that does not appear in the relations of $\mathcal{Q}$. Use Υ to perform a $(\mathcal{G}, \bigvee_{i=1}^m \mathcal{Z}_i)$ -sampling operation and simply return the output of this operation.

6 Random Enumeration from Sampling

This section will establish a connection between random enumeration and uniform sampling and then leverage the connection to prove Theorem 9.

Reduction. We consider a general setup. Let S be a set of elements whose size $|S|$ is known. Two subroutines are at our disposal:

- The first lists the elements of S (in an order we cannot control) within $|S| \cdot \lambda_{rep}$ time;
 - The second samples a uniformly random element of S with replacement within λ_{sam} time.
- We will show that the elements of S can be randomly enumerated with an expected delay of $O(\lambda_{rep} + \lambda_{sam})$.

Let us first describe an algorithm that produces a random permutation of S but does not ensure a small delay. The algorithm runs in three phases. In the first one, we use dynamic perfect hashing [20] to maintain a set S_{seen} of elements that have been found. Initially, $S_{seen} = \emptyset$. Then, we carry out *iterations*, each of which adds a new element to S_{seen} . Specifically, an iteration starts by randomly sampling an element e from S . If $e \notin S_{seen}$, we add it to S_{seen} and output it. Otherwise, the iteration re-samples from S until getting an unseen element. As long as $|S_{seen}| \leq |S|/2$, an unseen element is sampled with probability at least $1/2$. Hence, two samples suffice in expectation, and the cost of an iteration is $O(\lambda_{sam})$ expected. The first phase finishes after $|S|/2$ iterations. The second phase finds the entire S , and extracts (with hashing) and randomly permutes (with Fisher-Yates shuffle [22]) $S \setminus S_{seen}$. The phase performs at most $c \cdot |S| \cdot \lambda_{rep}$ atomic instructions of the RAM model for some constant c . Finally, the third phase outputs $S \setminus S_{seen}$ by the permuted order.

We apply a *de-amortization* approach to implement the above algorithm with a small expected delay. In the first phase, every time an element is added to S_{seen} , we do not output it immediately, but instead append it to a buffer queue S_{buf} . We remove and output the head element of S_{buf} every

$$\alpha = \left\lceil 1 + \frac{\lambda_{rep}}{\lambda_{sam}} \right\rceil$$

iterations so that S_{buf} still has $\frac{|S|}{2}(1 - 1/\alpha)$ elements left at the end of the first phase. In the second phase, we remove and enumerate the head of S_{buf} after every $c \cdot |S| \cdot \lambda_{rep} / (\frac{|S|}{2}(1 - \frac{1}{\alpha})) = 2c \cdot \lambda_{rep} \cdot \frac{\alpha}{\alpha-1}$ atomic instructions. The third phase simply enumerates the remaining elements in S_{buf} (by their queued order) and $S \setminus S_{seen}$ (by their permuted order) with a constant delay. Overall, the expected delay, which is determined by the first two phases, is bounded by $\alpha \cdot O(\lambda_{sam}) + \frac{2c \cdot \lambda_{rep} \cdot \alpha}{\alpha-1}$. It is rudimentary to show that the bound is $O(\lambda_{sam} + \lambda_{rep})$.

Proof of Theorem 9. Consider an input $\mathcal{Q}$ to the $(\mathcal{G}, \mathcal{Z})$ -sampling problem with a feasible pair $(\mathcal{G}, \mathcal{Z})$. Given a tuple $\mathbf{z}$ over $\mathcal{Z}$, define $S = \sigma_{\mathbf{z}}(\text{Join}(\mathcal{Q}))$. Our structure in Section 4.2 can be used to obtain the size $|S|$ in $O(1)$ time. Theorem 5 ensures $\lambda_{sam} = O(1)$. As explained in Appendix D, it is possible to build a structure that uses $O(\text{IN})$ space, can be built in $O(\text{IN})$ expected time, and can report S in $O(1 + |S|)$ time, implying $\lambda_{rep} = O(1)$. Using the reduction described earlier, we can randomly permute $\sigma_{\mathbf{z}}(\text{Join}(\mathcal{Q}))$ with an $O(1)$ expected delay. This completes the proof of Theorem 9. $\blacktriangleleft$

References

- 1 Serge Abiteboul, Richard Hull, and Victor Vianu. *Foundations of Databases*. Addison-Wesley, 1995.
- 2 Swarup Acharya, Phillip B. Gibbons, Viswanath Poosala, and Sridhar Ramaswamy. Join synopses for approximate query answering. In *SIGMOD*, pages 275–286, 1999. doi:10.1145/304182.304207.
- 3 Albert Atserias, Martin Grohe, and Daniel Marx. Size bounds and query plans for relational joins. *SIAM J. of Comp.*, 42(4):1737–1767, 2013. doi:10.1137/110859440.
- 4 Martin Aumuller, Sarel Har-Peled, Sepideh Mahabadi, Rasmus Pagh, and Francesco Silvestri. Sampling a near neighbor in high dimensions – Who is the fairest of them all? *TODS*, 47(1):4:1–4:40, 2022. doi:10.1145/3502867.
- 5 Guillaume Bagan, Arnaud Durand, and Etienne Grandjean. On acyclic conjunctive queries and constant delay enumeration. In *Computer Science Logic*, pages 208–222, 2007. doi:10.1007/978-3-540-74915-8_18.
- 6 Guillaume Bagan, Arnaud Durand, Etienne Grandjean, and Frederic Olive. Computing the j th solution of a first-order query. *RAIRO Theor. Informatics Appl.*, 42(1):147–164, 2008. doi:10.1051/ITA:2007046.
- 7 Karl Bringmann, Nofar Carmeli, and Stefan Mengel. Tight fine-grained bounds for direct access on join queries. *TODS*, 50(1):1:1–1:44, 2025. doi:10.1145/3707448.
- 8 Florent Capelli, Oliver Irwin, and Sylvain Salvati. A simple algorithm for worst case optimal join and sampling. In *ICDT*, pages 23:1–23:19, 2025. doi:10.4230/LIPIcs.ICDT.2025.23.
- 9 Nofar Carmeli and Markus Kroll. Enumeration complexity of conjunctive queries with functional dependencies. *Theory Comput. Syst.*, 64(5):828–860, 2020. doi:10.1007/S00224-019-09937-9.
- 10 Nofar Carmeli and Markus Kroll. On the enumeration complexity of unions of conjunctive queries. *TODS*, 46(2):5:1–5:41, 2021. doi:10.1145/3450263.
- 11 Nofar Carmeli and Luc Segoufin. Conjunctive queries with self-joins, towards a fine-grained enumeration complexity analysis. In *PODS*, pages 277–289, 2023. doi:10.1145/3584372.3588667.
- 12 Nofar Carmeli, Nikolaos Tziavelis, Wolfgang Gatterbauer, Benny Kimelfeld, and Mirek Riedewald. Tractable orders for direct access to ranked answers of conjunctive queries. *TODS*, 48(1):1:1–1:45, 2023. doi:10.1145/3578517.
- 13 Nofar Carmeli, Shai Zeevi, Christoph Berkholz, Alessio Conte, Benny Kimelfeld, and Nicole Schweikardt. Answering (unions of) conjunctive queries using random access and random-order enumeration. *TODS*, 47(3):9:1–9:49, 2022. doi:10.1145/3531055.

- 14 Timothy M. Chan, Saladi Rahul, and Jie Xue. Range closest-pair search in higher dimensions. *Computational Geometry*, 91:101669, 2020. doi:10.1016/J.COMGEO.2020.101669.
- 15 Surajit Chaudhuri, Rajeev Motwani, and Vivek R. Narasayya. On random sampling over joins. In *SIGMOD*, pages 263–274, 1999. doi:10.1145/304182.304206.
- 16 Yu Chen and Ke Yi. Random sampling and size estimation over cyclic joins. In *ICDT*, pages 7:1–7:18, 2020. doi:10.4230/LIPIcs.ICDT.2020.7.
- 17 Thomas H. Cormen, Charles E. Leiserson, Ronald L. Rivest, and Clifford Stein. *Introduction to Algorithms, Second Edition*. The MIT Press, 2001.
- 18 Shaleen Deep and Paraschos Koutris. Compressed representations of conjunctive query results. In *PODS*, pages 307–322, 2018. doi:10.1145/3196959.3196979.
- 19 Shiyuan Deng, Shangqi Lu, and Yufei Tao. On join sampling and the hardness of combinatorial output-sensitive join algorithms. In *PODS*, pages 99–111, 2023. doi:10.1145/3584372.3588666.
- 20 Martin Dietzfelbinger, Anna R. Karlin, Kurt Mehlhorn, Friedhelm Meyer auf der Heide, Hans Rohnert, and Robert Endre Tarjan. Dynamic perfect hashing: Upper and lower bounds. *SIAM J. of Comp.*, 23(4):738–761, 1994. doi:10.1137/S0097539791194094.
- 21 Arnaud Durand. Fine-grained complexity analysis of queries: From decision to counting and enumeration. In *PODS*, pages 331–346, 2020. doi:10.1145/3375395.3389130.
- 22 Ronald Aylmer Fisher and Frank Yates. *Statistical Tables for Biological, Agricultural and Medical Research*. Edinburgh and London: Oliver and Boyd., 1953.
- 23 Isaac Goldstein, Tsvi Kopelowitz, Moshe Lewenstein, and Ely Porat. Conditional lower bounds for space/time tradeoffs. In *WADS*, pages 421–436. Springer, 2017. doi:10.1007/978-3-319-62127-2_36.
- 24 Isaac Goldstein, Moshe Lewenstein, and Ely Porat. On the hardness of set disjointness and set intersection with bounded universe. In *ISAAC*, pages 7:1–7:22, 2019. doi:10.4230/LIPIcs.ISAAC.2019.7.
- 25 Norman L Johnson, Adrienne W Kemp, and Samuel Kotz. *Univariate Discrete Distributions*. John Wiley & Sons, 2005.
- 26 Richard M. Karp, Michael Luby, and Neal Madras. Monte-carlo approximation algorithms for enumeration problems. *J. Algorithms*, 10(3):429–448, 1989. doi:10.1016/0196-6774(89)90038-2.
- 27 Mahmoud Abo Khamis, Hung Q. Ngo, and Dan Suciu. What do shannon-type inequalities, submodular width, and disjunctive datalog have to do with one another? In *PODS*, pages 429–444, 2017. doi:10.1145/3034786.3056105.
- 28 Kyoungmin Kim, Jaehyun Ha, George Fletcher, and Wook-Shin Han. Guaranteeing the $\tilde{O}(\text{AGM}/\text{OUT})$ runtime for uniform sampling and size estimation over joins. In *PODS*, pages 113–125, 2023. doi:10.1145/3584372.3588676.
- 29 Shangqi Lu and Yufei Tao. Indexing for keyword search with structured constraints. In *PODS*, pages 263–275, 2023. doi:10.1145/3584372.3588663.
- 30 Mihai Patrascu and Liam Roditty. Distance oracles beyond the Thorup-Zwick bound. *SIAM J. of Comp.*, 43(1):300–311, 2014. doi:10.1137/11084128X.
- 31 Alastair J. Walker. New fast method for generating discrete random numbers with arbitrary frequency distributions. *Electronics Letters*, 10(8):127–128, 1974.
- 32 Ru Wang and Yufei Tao. Join and subgraph sampling under degree constraints. *JCSS*, 155, 2026. doi:10.1016/J.JCSS.2025.103693.
- 33 Mihalis Yannakakis. Algorithms for acyclic database schemes. In *VLDB*, pages 82–94, 1981.
- 34 Hangdong Zhao, Shaleen Deep, and Paraschos Koutris. Space-time tradeoffs for conjunctive queries with access patterns. In *PODS*, pages 59–68, 2023. doi:10.1145/3584372.3588675.
- 35 Zhuoyue Zhao, Robert Christensen, Feifei Li, Xiao Hu, and Ke Yi. Random sampling over joins revisited. In *SIGMOD*, pages 1525–1539, 2018. doi:10.1145/3183713.3183739.

A

 Supplementary Proofs for Section 3.2

A.1 Proof of Lemma 10

We first review a property of acyclic joins that can be used to compute, for each tuple $\mathbf{u} \in R_e$, the result of $\mathbf{u} \bowtie \text{Join}(\mathcal{Q}_e)$ in a manner reminiscent of Cartesian products. W.l.o.g., suppose that e has $t \geq 1$ child nodes in $\mathcal{T}$ denoted as $e_1, e_2, \dots, e_t$. Define $\mathcal{X}_i = e \cap e_i$ for each $i \in [t]$. Consider the set $S_\times$ of tuples created by the procedure below:

children-prod ($\mathbf{u}$) $\text{ /* } \mathbf{u}$ is a tuple in $R_e \text{ */}$

1. $S_\times \leftarrow \emptyset$
2. $S_i \leftarrow \mathbf{u}[\mathcal{X}_i] \bowtie \text{Join}(\mathcal{Q}_{e_i})$ for each $i \in [t]$
 $\text{ /* recall that } \mathbf{u}[\mathcal{X}_i] \text{ is the projection of } \mathbf{u} \text{ onto } \mathcal{X}_i \text{ */}$
3. **for** each $(\mathbf{v}_1, \dots, \mathbf{v}_t) \in S_1 \times \dots \times S_t$ **do**
4. add $\mathbf{u} \bowtie \mathbf{v}_1 \bowtie \mathbf{v}_2 \bowtie \dots \bowtie \mathbf{v}_t$ to $S_\times$

The next lemma is a well-known property of Yannakakis' algorithm [33].

► **Lemma 22** ([1, 33]). *Both statements below are true regarding **children-prod**: (i) Line 4 always adds a new tuple to $S_\times$, and (ii) $\mathbf{u} \bowtie \text{Join}(\mathcal{Q}_e) = S_\times$.*

We are now ready to prove Lemma 10 by induction. If e (the node given in the lemma) is a leaf of the (rooted) $\mathcal{T}$, then $\mathcal{Q}_e$ includes only R_e . In this case, $\text{Join}(\mathcal{Q}_e) = R_e$ and $\text{subj-}w_e(\mathbf{u}) = |\mathbf{u} \bowtie R_e| = 1$. Hence, the lemma holds on R_e .

Now, consider e to be an internal node of $\mathcal{T}$. W.l.o.g., suppose that e has $t \geq 1$ child nodes in $\mathcal{T}$ denoted as $e_1, \dots, e_t$. Assuming inductively that the lemma is correct on R_{e_i} of every $i \in [t]$, next we will prove the correctness on R_e as well. For each $i \in [t]$, define $\mathcal{X}_i = e \cap e_i$.

Let $\mathbf{u}$ be an arbitrary tuple in R_e . According to Lemma 22, we have

$$\text{subj-}w_e(\mathbf{u}) = \prod_{i=1}^t |S_i(\mathbf{u})|$$

where $S_i(\mathbf{u}) = \mathbf{u}[\mathcal{X}_i] \bowtie \text{Join}(\mathcal{Q}_{e_i})$, as defined at Line 2 of **children-prod**. For each $i \in [t]$, we will show how to construct in $O(\text{IN})$ expected time a structure of $O(\text{IN})$ space from which we can obtain $|S_i(\mathbf{u})|$ for every $\mathbf{u} \in R_e$ in $O(\text{IN})$ time. This will prove Lemma 10 because $t = O(1)$.

For each $\mathbf{u} \in R_e$, define

$$R_{e_i}(\mathbf{u}) = \{\mathbf{v} \in R_{e_i} \mid \mathbf{v}[\mathcal{X}_i] = \mathbf{u}[\mathcal{X}_i]\}. \quad (11)$$

We can now write

$$\begin{aligned} |S_i(\mathbf{u})| &= |\mathbf{u}[\mathcal{X}_i] \bowtie \text{Join}(\mathcal{Q}_{e_i})| = \left| \bigcup_{\mathbf{v} \in R_{e_i}(\mathbf{u})} \mathbf{v} \bowtie \text{Join}(\mathcal{Q}_{e_i}) \right| \\ &= \sum_{\mathbf{v} \in R_{e_i}(\mathbf{u})} |\mathbf{v} \bowtie \text{Join}(\mathcal{Q}_{e_i})| = \sum_{\mathbf{v} \in R_{e_i}(\mathbf{u})} \text{subj-}w_{e_i}(\mathbf{v}). \end{aligned} \quad (12)$$

We can obtain $|S_i(\mathbf{u})|$ with one scan over $R_{e_i}(\mathbf{u})$ because the subjoin weights of the tuples in R_{e_i} have been computed (by induction). To identify $R_{e_i}(\mathbf{u})$ for $\mathbf{u}$, we create a perfect-hashing structure on R_{e_i} in $O(|R_{e_i}|)$ expected time, after which $R_{e_i}(\mathbf{u})$ can be retrieved in $O(1 + |R_{e_i}(\mathbf{u})|)$ time.

We divide R_e into equivalent classes based on the tuples' projections onto $\mathcal{X}_i$. If two tuples $\mathbf{u}_1$ and $\mathbf{u}_2$ are in the same equivalent class, the numbers $|S_i(\mathbf{u}_1)|$ and $|S_i(\mathbf{u}_2)|$ are identical. Hence, it suffices to compute $|S_i(\mathbf{u})|$ for one (arbitrary) representative $\mathbf{u}$ of each equivalent class. On the other hand, if $\mathbf{u}_1$ and $\mathbf{u}_2$ are from different equivalent classes, the sets $R_{e_i}(\mathbf{u}_1)$ and $R_{e_i}(\mathbf{u}_2)$ are disjoint. Hence, we can compute the $|S_i(\mathbf{u})|$ values for all representatives in time

$$\sum_{\text{representative } \mathbf{u} \text{ of } R_e} O(1 + |R_{e_i}(\mathbf{u})|) = O(|R_e| + |R_{e_i}|) = O(\text{IN}).$$

A.2 Proof of Lemma 11

We prove the lemma by induction. The base case where e is a leaf of $\mathcal{T}$ is trivial and omitted.

Now, consider e to be an internal node of $\mathcal{T}$. W.l.o.g., suppose that e has $t \geq 1$ child nodes in $\mathcal{T}$ denoted as $e_1, \dots, e_t$. Assuming inductively that the lemma is correct on R_{e_i} of every $i \in [t]$, next we will prove the correctness on R_e as well. For each $i \in [t]$, define $\mathcal{X}_i = e \cap e_i$.

Let $\mathbf{u}$ be an arbitrary tuple in R_e . According to Lemma 22, there is a one-one correspondence between $\mathbf{u} \bowtie \text{Join}(\mathcal{Q}_e)$ and $S_1(\mathbf{u}) \times S_2(\mathbf{u}) \times \dots \times S_t(\mathbf{u})$ where $S_i(\mathbf{u}) = \mathbf{u}[\mathcal{X}_i] \bowtie \text{Join}(\mathcal{Q}_{e_i})$, as defined at Line 2 of procedure **children-prod**. Hence, to take a uniform random tuple from $\mathbf{u} \bowtie \text{Join}(\mathcal{Q}_e)$, we can take a uniformly random tuple from $S_i(\mathbf{u})$ for each $i \in [t]$. Next, we will show how to construct in $O(\text{IN})$ expected time a structure of $O(\text{IN})$ space from which we can sample from $S_i(\mathbf{u})$ in $O(1)$ time. This will prove Lemma 11 because $t = O(1)$.

For each $\mathbf{u} \in R_e$, define $R_{e_i}(\mathbf{u})$ as in (11). We can now write

$$S_i(\mathbf{u}) = \mathbf{u}[\mathcal{X}_i] \bowtie \text{Join}(\mathcal{Q}_{e_i}) = \bigcup_{\mathbf{v} \in R_{e_i}(\mathbf{u})} \mathbf{v} \bowtie \text{Join}(\mathcal{Q}_{e_i}).$$

Random sampling from $S_i(\mathbf{u})$ can be performed in two steps. First, obtain a random tuple $\mathbf{X} \in R_e(\mathbf{u})$ such that $\Pr[\mathbf{X} = \mathbf{v}] = |\mathbf{v} \bowtie \text{Join}(\mathcal{Q}_{e_i})| / |S_i(\mathbf{u})| = \text{subj-}w_{e_i}(\mathbf{v}) / |S_i(\mathbf{u})|$ for each $\mathbf{v} \in R_{e_i}(\mathbf{u})$. Second, draw a uniformly random tuple from $\mathbf{X} \bowtie \text{Join}(\mathcal{Q}_{e_i})$.

Since the second step can be supported in $O(1)$ time by induction, we focus on the first step. As $|S_i(\mathbf{u})| = \sum_{\mathbf{v} \in R_{e_i}(\mathbf{u})} \text{subj-}w_{e_i}(\mathbf{v})$, the first step is an instance of weighted sampling once the $\text{subj-}w_{e_i}(\mathbf{v})$ of every $\mathbf{v} \in R_{e_i}$ has been computed using Lemma 10 in $O(\text{IN})$ expected time. An alias structure on R_{e_i} allows us to obtain $\mathbf{X}$ in constant time. The structure can be built in $O(1 + |R_{e_i}(\mathbf{u})|)$ time. It suffices to build such a structure for every “equivalent class” of R_e (as defined in Section A.1). All those structures take $O(\text{IN})$ time to construct in total and occupy $O(\text{IN})$ space overall.

B Correctness of Our Sampling Algorithm in Case 2 of Section 4.2

Consider the following procedure:

component-prod

1. $S_\times \leftarrow \emptyset$
2. $S_i \leftarrow \sigma_{\mathbf{z}_i}(\text{Join}(\mathcal{Q}_i))$ for each $i \in [s]$ /* recall that $\mathbf{z}_i = \mathbf{z}[\mathcal{V}_i]$ */
3. **for** each $(\mathbf{v}_1, \dots, \mathbf{v}_s) \in S_1 \times \dots \times S_s$ **do**
4. add $\mathbf{v}_1 \bowtie \mathbf{v}_2 \bowtie \dots \bowtie \mathbf{v}_s$ to $S_\times$

The correctness of our sampling algorithm for Case 2 of Section 4.2 is a corollary of:

► **Lemma 23.** *Both statements below are true about **component-prod**: (i) Line 4 always adds a new tuple to $S_\times$, and (ii) $\sigma_{\mathbf{z}}(\text{Join}(\mathcal{Q})) = S_\times$.*

Proof.

Proof of Statement (i). For each $i \in [s]$, take an arbitrary tuple $\mathbf{v}_i \in S_i$. We will first show $\mathbf{v}_1 \bowtie \mathbf{v}_2 \bowtie \dots \bowtie \mathbf{v}_s \neq \emptyset$. For this purpose, it suffices to prove that, for any distinct $i, j \in [s]$ with $\mathcal{V}_i \cap \mathcal{V}_j \neq \emptyset$, it must hold that $\mathbf{v}_i[Y] = \mathbf{v}_j[Y]$ for any attribute $Y \in \mathcal{V}_i \cap \mathcal{V}_j$.

As $Y \in \mathcal{V}_i$ (resp., $Y \in \mathcal{V}_j$), there is a hyperedge $e_i \in \mathcal{E}_i$ (resp., $e_j \in \mathcal{E}_j$) containing Y . As e_i and e_j are in different $(\mathcal{Z}, \mathcal{T})$ -components – recall that $\mathcal{T}$ is a join tree of $\mathcal{Q}$ – the (unique) simple path between them on $\mathcal{T}$ must cross at least one $\mathcal{Z}$ -breakable edge, denoted as $\{e_\circ, e_\bullet\}$. The connectedness property of $\mathcal{T}$ ensures that $Y \in e_\circ \cap e_\bullet$. By the definition of $\mathcal{Z}$ -breakable edge, Y must be an attribute in $\mathcal{Z}$, indicating that the tuple $\mathbf{z}$ has a Y -value. As $\mathbf{v}_i \in S_i$ and $\mathbf{v}_j \in S_j$, we must have $\mathbf{v}_i(Y) = \mathbf{v}_j(Y) = \mathbf{z}(Y)$.

To prove the statement, it remains to show that Line 4 never adds to $S_\times$ the same tuple twice. Consider any two executions of Line 4: the first with $\mathbf{v}_i = \mathbf{v}_i^1$ for $i \in [s]$, and the second with $\mathbf{v}_i = \mathbf{v}_i^2$ for $i \in [s]$. There exists at least one $j \in [s]$ such that $\mathbf{v}_j^1 \neq \mathbf{v}_j^2$. This further indicates the existence of an attribute $Y \in \mathcal{V}_j$ such that $\mathbf{v}_j^1(Y) \neq \mathbf{v}_j^2(Y)$. We can now assert that $\mathbf{v}_1^1 \bowtie \dots \bowtie \mathbf{v}_s^1$ and $\mathbf{v}_1^2 \bowtie \dots \bowtie \mathbf{v}_s^2$ must differ in their Y -values.

Proof of Statement (ii). We will first prove

$$\sigma_{\mathbf{z}}(\text{Join}(\mathcal{Q})) \subseteq S_\times. \quad (13)$$

Take any tuple $\mathbf{v} \in \sigma_{\mathbf{z}}(\text{Join}(\mathcal{Q}))$. Define $\mathbf{v}_i = \mathbf{v}[\mathcal{V}_i]$ for each $i \in [s]$. To show (13), it suffices to prove $\mathbf{v}_i \in \sigma_{\mathbf{z}_i}(\text{Join}(\mathcal{Q}_i))$ for every $i \in [s]$. We achieve the purpose by arguing that $\mathbf{v}_i[e] \in R_e$ for any $e \in \mathcal{E}_i$. As every relation of $\mathcal{Q}_i$ belongs to $\mathcal{Q}$, we know $R_e \in \mathcal{Q}$. We can thus infer $\mathbf{v}[e] \in R_e$ from $\mathbf{v} \in \sigma_{\mathbf{z}}(\text{Join}(\mathcal{Q}))$. On the other hand, $\mathbf{v}[e] = \mathbf{v}_i[e]$ because $e \subseteq \mathcal{V}_i$. It thus follows that $\mathbf{v}_i[e] \in R_e$.

Next, we will prove

$$S_\times \subseteq \sigma_{\mathbf{z}}(\text{Join}(\mathcal{Q})). \quad (14)$$

For each $i \in [s]$, take an arbitrary tuple $\mathbf{v}_i \in S_i$. Define $\mathbf{w} = \mathbf{v}_1 \bowtie \dots \bowtie \mathbf{v}_s$. Our proof of statement (i) has explained that $\mathbf{w}$ cannot be empty. To prove (14), our goal is to show that $\mathbf{w} \in \sigma_{\mathbf{z}}(\text{Join}(\mathcal{Q}))$. We first argue that $\mathbf{w} \in \text{Join}(\mathcal{Q})$, i.e., $\mathbf{w}[e] \in R_e$ for every $e \in \mathcal{E}$. Indeed, as e must belong to $\mathcal{E}_i$ for some $i \in [s]$, we know $\mathbf{w}[e] = \mathbf{v}_i[e]$, which is in R_e because $\mathbf{v}_i \in \sigma_{\mathbf{z}_i}(\text{Join}(\mathcal{Q}_i))$. To prove $\mathbf{w} \in \sigma_{\mathbf{z}}(\text{Join}(\mathcal{Q}))$, we still have to show $\mathbf{w}(Z) = \mathbf{z}(Z)$ for every $Z \in \mathcal{Z}$. For this purpose, simply identify any $i \in [s]$ satisfying $Z \in \mathcal{V}_i$ (this i exists because $\bigcup_{i=1}^s \mathcal{V}_i = \mathcal{V}$). That $\mathbf{w}(Z) = \mathbf{z}(Z)$ follows from the fact $\mathbf{v}_i \in \sigma_{\mathbf{z}_i}(\text{Join}(\mathcal{Q}_i))$.

Statement (ii) follows from (13) and (14). $\blacktriangleleft$

C Supplementary Content for Section 5

Proof of Proposition 19. For any random variables X and Y , the total law of expectation states $\mathbf{E}[X] = \mathbf{E}[\mathbf{E}[X \mid Y]]$. To prove the proposition, define X to be 1 if $U \leq \Gamma$ or 0 otherwise; furthermore, set $Y = \Gamma$. Clearly, $\mathbf{E}[X \mid Y] = \Pr[U \leq \Gamma \mid \Gamma] = \Gamma/m$ by the independence of U and Γ . We can now derive $\Pr[U \leq \Gamma] = \mathbf{E}[X] = \mathbf{E}[\mathbf{E}[X \mid Y]] = \mathbf{E}[\Gamma/m] = \frac{1}{m} \mathbf{E}[\Gamma]$. $\blacktriangleleft$

Proof of Proposition 20. If $\deg(y) = 1$, then F is deterministically $m - 1$, in which case the lemma clearly holds.

The rest of the proof considers $\deg(y) \geq 2$. In that scenario, F can be rephrased in the following conventional WoR-sampling setup. Suppose that we have $m - 1$ balls, among which $\deg(y) - 1$ ones are red and the rest are white. Uniformly sample a ball WoR until seeing a red ball, and F gives the number of white balls sampled in this process. It is well known (e.g., see [25]) that F follows the Negative Hypergeometric distribution with $\mathbf{E}[F] = \frac{m}{\deg(y)} - 1$. $\blacktriangleleft$

Supporting the Size Operation in Section 5.2. Let $\mathcal{G} = (\mathcal{V}, \mathcal{E})$ be a schema graph and $\mathcal{Z}$ be a subset of $\mathcal{V}$ such that $\mathcal{G}$ is ext- $\mathcal{Z}$ -connex. Let $\mathcal{Q}$ be a join instance of $\mathcal{G}$. Suppose that we have built a structure Υ of Theorem 5 on $\mathcal{Q}$ based on the description in Section 4.2. Given any tuple $\mathbf{z}$ over $\mathcal{Z}$, we will show how to use Υ to get the size $|\sigma_{\mathbf{z}}(\text{Join}(\mathcal{Q}))|$ in constant time. In fact, we have already discussed this in the scenario where $\mathcal{G}$ is $\mathcal{Z}$ -canonical – see the remark in Section 4.2. Next, we consider that $\mathcal{G}$ is not $\mathcal{Z}$ -canonical.

As $\mathcal{G}$ is ext- $\mathcal{Z}$ -connex, by Lemma 16, any join tree $\mathcal{T}$ of $\mathcal{G}$ defines a number s of $(\mathcal{Z}, \mathcal{T})$ -components of $\mathcal{G}$ – denoted as $\mathcal{G}_1 = (\mathcal{V}_1, \mathcal{E}_1), \dots, \mathcal{G}_s = (\mathcal{V}_s, \mathcal{E}_s)$, respectively – such that $\mathcal{G}_i$ is $(\mathcal{V}_i \cap \mathcal{Z})$ -canonical for every $i \in [s]$. For each $i \in [s]$, define $\mathcal{Z}_i = \mathcal{V}_i \cap \mathcal{Z}$ and $\mathcal{Q}_i = \{R_e \in \mathcal{Q} \mid e \in \mathcal{E}_i\}$. Note that $\mathcal{Q}_i$ is a join instance of $\mathcal{G}_i$, which is $\mathcal{Z}_i$ -canonical. For each $i \in [s]$, the structure Υ includes a structure of Theorem 5 on every $\mathcal{Q}_i$ ($i \in [s]$); denote that structure as Υ_i . Now, consider a tuple $\mathbf{z}$ over $\mathcal{Z}$ whose $|\sigma_{\mathbf{z}}(\text{Join}(\mathcal{Q}))|$ is to be computed. For each $i \in [s]$, we define $\mathbf{z}_i = \mathbf{z}[\mathcal{V}_i]$ and use Υ_i to obtain the size $|\sigma_{\mathbf{z}_i}(\text{Join}(\mathcal{Q}_i))|$ in $O(1)$ time (this is doable because $\mathcal{G}_i$ is $\mathcal{Z}_i$ -canonical). Finally, we return $|\sigma_{\mathbf{z}}(\text{Join}(\mathcal{Q}))| = \prod_{i=1}^s |\sigma_{\mathbf{z}_i}(\text{Join}(\mathcal{Q}_i))|$. The correctness is guaranteed by Lemma 23.

D Completing the Proof of Theorem 9

Let $\mathcal{G} = (\mathcal{V}, \mathcal{E})$ be an acyclic schema graph and $\mathcal{Q}$ be a join instance of $\mathcal{G}$. Yannakakis’ algorithm [33] outputs $\text{Join}(\mathcal{Q})$ in $O(\text{IN} + |\text{Join}(\mathcal{Q})|)$ time. The algorithm allows one to specify a *source relation* $R \in \mathcal{Q}$. In an $O(\text{IN})$ -time preprocessing stage, it

- converts R into a *fully-reduced* state, i.e., eliminating every tuple $\mathbf{u} \in R$ with $\mathbf{u} \bowtie \text{Join}(\mathcal{Q}) = \emptyset$ (i.e., $\mathbf{u}$ does not “contribute” to the join result);
- creates a structure that, given any tuple $\mathbf{u} \in R$, can report $\mathbf{u} \bowtie \text{Join}(\mathcal{Q})$ in $O(|\mathbf{u} \bowtie \text{Join}(\mathcal{Q})|)$ time.

Let $\mathcal{Z}$ be a subset of $\mathcal{V}$ such that $\mathcal{G}$ is ext- $\mathcal{Z}$ -connex. It should have become fairly straightforward to combine the above and the discussion in Section 5.2 to obtain a structure with the following guarantees: it uses $O(\text{IN})$ space, can be built in $O(\text{IN})$ expected time, and when given a tuple $\mathbf{z}$ over $\mathcal{Z}$, can report $\sigma_{\mathbf{z}}(\text{Join}(\mathcal{Q}))$ in $O(1 + |\sigma_{\mathbf{z}}(\text{Join}(\mathcal{Q}))|)$ time.

Case 1: $\mathcal{G}$ is $\mathcal{Z}$ -Canonical. By Definition 12, there is a hyperedge $e^* \in \mathcal{E}$ satisfying $\mathcal{Z} \subseteq e^*$. Run the preprocessing stage of Yannakakis’s algorithm to create the aforementioned structure Υ by specifying R_{e^*} as the source relation (recall that this is the relation in $\mathcal{Q}$ whose schema is e^*). The preprocessing leaves R_{e^*} in a fully-reduced state. Given a tuple $\mathbf{z}$ over $\mathcal{Z}$, we identify with hashing in constant time the set $R_{e^*}(\mathbf{z})$ as defined in (7). Then, for each $\mathbf{u} \in R_{e^*}(\mathbf{z})$, use Υ to report $\mathbf{u} \bowtie \text{Join}(\mathcal{Q})$ (which must be non-empty).

Case 2: $\mathcal{G}$ Is Not $\mathcal{Z}$ -Canonical. Let $\mathcal{T}$ be an arbitrary join tree $\mathcal{T}$ of $\mathcal{G}$, which defines a number s of $(\mathcal{Z}, \mathcal{T})$ -components of $\mathcal{G}$ – denoted as $\mathcal{G}_1 = (\mathcal{V}_1, \mathcal{E}_1), \dots, \mathcal{G}_s = (\mathcal{V}_s, \mathcal{E}_s)$, respectively – such that $\mathcal{G}_i$ is $(\mathcal{V}_i \cap \mathcal{Z})$ -canonical for every $i \in [s]$ (see Lemma 16). For each $i \in [s]$, define $\mathcal{Z}_i = \mathcal{V}_i \cap \mathcal{Z}$ and $\mathcal{Q}_i = \{R_e \in \mathcal{Q} \mid e \in \mathcal{E}_i\}$. Note that $\mathcal{Q}_i$ is a join instance of the schema graph $\mathcal{G}_i$, which is $\mathcal{Z}_i$ -canonical. For each $i \in [s]$, build a structure Υ_i on $\mathcal{Q}_i$ in the way explained for Case 1. Consider now a tuple $\mathbf{z}$ over $\mathcal{Z}$. For each $i \in [s]$, defining $\mathbf{z}_i = \mathbf{z}[\mathcal{V}_i]$, we use Υ_i to extract $S_i = \sigma_{\mathbf{z}_i}(\text{Join}(\mathcal{Q}_i))$. Then, for every $(\mathbf{u}_1, \dots, \mathbf{u}_s) \in S_1 \times \dots \times S_s$, output $\mathbf{u}_1 \bowtie \mathbf{u}_2 \bowtie \dots \bowtie \mathbf{u}_s$.