

Semi-Oblivious Chase Termination for Linear Existential Rules: An Experimental Study

Marco Calautti
University of Milan
marco.calautti@unimi.it

Mostafa Milani
University of Western Ontario
mostafa.milani@uwo.ca

Andreas Pieris
University of Edinburgh &
University of Cyprus
apieris@inf.ed.ac.uk

ABSTRACT

The chase procedure is a fundamental algorithmic tool in databases that allows us to reason with constraints, such as existential rules, with a plethora of applications. It takes as input a database and a set of constraints, and iteratively completes the database as dictated by the constraints. A key challenge, though, is the fact that it may not terminate, which leads to the problem of checking whether it terminates given a database and a set of constraints. In this work, we focus on the semi-oblivious version of the chase, which is well-suited for practical implementations, and linear existential rules, a central class of constraints with several applications. In this setting, there is a mature body of theoretical work that provides syntactic characterizations of when the chase terminates, algorithms for checking chase termination, and precise complexity results. Our main objective is to experimentally evaluate the existing chase termination algorithms with the aim of understanding which input parameters affect their performance, clarifying whether they can be used in practice, and revealing their performance limitations.

PVLDB Reference Format:

Marco Calautti, Mostafa Milani, and Andreas Pieris. Semi-Oblivious Chase Termination for Linear Existential Rules: An Experimental Study. PVLDB, 16(11): 2858 - 2870, 2023.
doi:10.14778/3611479.3611493

PVLDB Artifact Availability:

The source code, data, and/or other artifacts have been made available at <https://github.com/mostafamilani/chase-termination>.

1 INTRODUCTION

The *chase procedure* (or simply chase) is a fundamental algorithmic tool that has been successfully applied to several database problems such as checking logical implication of constraints [3, 15], containment of queries under constraints [1], computing data exchange solutions [10], and ontological query answering [7], to name a few. The chase takes as input a database D and a set Σ of constraints, which, for this work, are *existential rules* (a.k.a. *tuple-generating dependencies* (TGDs)) of the form $\forall \bar{x} \forall \bar{y} (\phi(\bar{x}, \bar{y}) \rightarrow \exists \bar{z} \psi(\bar{x}, \bar{z}))$, where ϕ (the body) and ψ (the head) are conjunctions of relational atoms, and it produces an instance D_Σ that is a *universal model* of D and Σ , i.e., a model that can be homomorphically embedded into every

other model of D and Σ . Somehow D_Σ acts as a representative of all the models of D and Σ . This is the reason for the ubiquity of the chase in databases, as discussed in [9]. Indeed, many database problems can be solved by simply exhibiting a universal model.

Roughly speaking, the chase adds new tuples to the database D (possibly with null values that act as witnesses for the existentially quantified variables), as dictated by the TGDs of Σ , and it keeps doing this until all the TGDs of Σ are satisfied. There are, in principle, two practically relevant ways for formalizing this simple idea, which lead to different versions of the chase: *semi-oblivious* and *restricted*. The key difference between those two versions of the chase is when a TGD is considered applicable. In a nutshell, the semi-oblivious version applies a TGD whenever the body is satisfied, whereas the restricted version applies a TGD if the body is satisfied but the head is not. Here is a simple example that illustrates the differences between the two versions of the chase procedure; for a more detailed discussion, we refer the reader to [6].

Example 1.1. Consider the database $D = \{R(a, a)\}$ and the TGD $\forall x \forall y (R(x, y) \rightarrow \exists z R(z, x))$. Although $R(a, a)$ satisfies the body of the TGD, the restricted chase will detect that the database already satisfies the TGD, and thus, there is nothing to derive. On the other hand, the semi-oblivious chase will produce the atom $R(\perp_1, a)$, where $\perp_1$ is a (labelled) null value, despite the fact that the TGD is already satisfied. Observe now that $R(\perp_1, a)$ satisfies the body of the TGD, and the semi-oblivious chase will produce the atom $R(\perp_2, \perp_1)$, where $\perp_2$ is a null value. The semi-oblivious chase will keep indefinitely applying the TGD, and eventually will build the infinite instance $\{R(a, a), R(\perp_1, a), R(\perp_2, \perp_1), R(\perp_3, \perp_2), \dots\}$. ■

The restricted version of the chase has a clear advantage over the semi-oblivious one as it generally builds smaller instances. But, of course, this advantage does not come for free: at each step, the restricted chase has to check that there is no way to satisfy the head of the TGD at hand, and this can be very costly in practice. It has been recently observed that for RAM-based implementations the restricted chase is the indicated approach since the benefit from producing smaller instances justifies the additional effort for checking whether a TGD is already satisfied; see, e.g., [4, 13]. However, as discussed in [4], an RDBMS-based implementation of the restricted chase is quite challenging, whereas an efficient implementation of the semi-oblivious chase is feasible. Hence, both the semi-oblivious and restricted versions of the chase are relevant tools for practical implementations.

Chase Termination and Linear TGDs. There are indeed efficient implementations of the semi-oblivious and restricted chase that allow us to solve important database problems by adopting a materialization-based approach [4, 13, 17, 19]. Nevertheless, for

This work is licensed under the Creative Commons BY-NC-ND 4.0 International License. Visit <https://creativecommons.org/licenses/by-nc-nd/4.0/> to view a copy of this license. For any use beyond those covered by this license, obtain permission by emailing info@vldb.org. Copyright is held by the owner/author(s). Publication rights licensed to the VLDB Endowment.
Proceedings of the VLDB Endowment, Vol. 16, No. 11 ISSN 2150-8097.
doi:10.14778/3611479.3611493

this to be feasible in practice we need a guarantee that the chase terminates. There are settings where the termination of the chase is guaranteed by design. For example, in data exchange there is an a priori assumption that the chase terminates as the adopted languages (e.g., weakly-acyclic TGDs) are deliberately designed to ensure the termination of the chase [10]. There are settings, however, where such an a priori assumption cannot be realistically made. Such a setting, which is of special interest for our work, is that of ontological reasoning. Indeed, even lightweight TGD-based ontology languages allow us to express non-acyclic statements with value invention, a combination that may lead to an infinite chase; see, e.g., [11]. This fact motivated a long line of research on the chase termination problem, that is, given a database D and a set Σ of TGDs, to devise sound and complete algorithms that check whether the semi-oblivious or restricted chase of D with Σ terminates. It is known that, in general, this is an undecidable problem. This has been established in [9] for the restricted chase, and it was observed a year later in [16] that the same proof shows undecidability also for the semi-oblivious chase. The undecidability proof given in [9], however, constructs a sophisticated set of TGDs that goes beyond existing well-behaved classes of TGDs that enjoy certain syntactic properties. This observation leads to the obvious question: is the chase termination problem algorithmically solvable whenever we focus on well-behaved classes of TGDs?

A well-behaved class of TGDs, which forms a central ontology language that attracted a considerable attention due to its simplicity and the fact that it strikes a good balance between expressiveness and complexity, is that of linear TGDs proposed in [7]. A TGD is *linear* if it has only one atom in its body, whereas the head can be an arbitrary conjunction of atoms. Such a TGD is called *simple-linear* if each variable in its body occurs only once, whereas variables in its head can repeat without any restriction. Although, at first glance, (simple-)linear TGDs may look very inexpressive, it turns out that they are powerful enough for modelling ontologies. In particular, the practically important ontology language DL-Lite_R [8], which is based on Description Logics and forms the logical underpinning of OWL 2 QL, one of the popular profiles of the W3C committee's Web Ontology Language (OWL) standard for ontology languages, can be easily embedded into the class of simple-linear TGDs.

The chase termination problem in the presence of (simple-)linear TGDs has been extensively studied the last few years. Concerning the semi-oblivious version of the chase, there is a mature body of theoretical work that provides syntactic characterizations of when the chase terminates based on suitable acyclicity notions, algorithms for checking chase termination, precise complexity results, and worst-case optimal bounds on the size of the chase instance (whenever is finite) [5]. On the other hand, for the restricted version of the chase, we only have a decidability result via an algorithm that runs in at least triple-exponential time, under the assumption that the head of the linear TGDs consists of a single atom; in fact, this algorithm relies on an expensive reduction to the satisfiability problem of Monadic Second Order Logic over infinite trees [14]. Needless to say that the syntactic characterizations of the termination of the semi-oblivious chase, as well as the existing algorithms for checking the termination of the semi-oblivious chase, are far from being applicable to the restricted chase, and there is no way to be merely adapted in order to be applicable to the restricted chase.

This striking difference on the progress that has been achieved should be attributed to the fact that the chase termination problem is significantly more challenging in the case of the restricted chase.

Main Objective. Having a complete theoretical understanding of the semi-oblivious chase termination problem in the presence of (simple-)linear TGDs, the next step is to experimentally evaluate the existing algorithms with the aim of understanding which input parameters affect their performance, clarifying whether they can be applied in a practical context, and revealing their performance limitations. This is the main objective of this work. We do not consider the restricted chase as this will be premature due to the lack of a good theoretical understanding of the problem as discussed above; in fact, this is the subject of an ongoing research activity.

From the chase termination literature, we can inherit two types of algorithms for the semi-oblivious chase termination problem in the presence of (simple-)linear TGDs, that is, *materialization-based* and *acyclicity-based* [5], which can be described as follows.

The **materialization-based** algorithms exploit the existence of worst-case optimal bounds on the size of the chase instance (whenever is finite). In particular, given a database D and a set Σ of (simple-)linear TGDs, we have an integer $k_{D,\Sigma}$ such that the chase of D with Σ terminates iff the size of the chase instance (i.e., the number of its atoms) is at most $k_{D,\Sigma}$. This immediately leads to conceptually simple chase termination algorithms: simply run the semi-oblivious chase of D with Σ and keep a counter for the number of generated atoms, and if the count exceeds $k_{D,\Sigma}$, then conclude that the chase does not terminate; otherwise, it does.

On the other hand, the **acyclicity-based** algorithms exploit the syntactic characterizations of when the chase terminates via suitable acyclicity notions. In particular, given a database D and a set Σ of (simple-)linear TGDs, we know that the chase of D with Σ terminates iff the dependency graph of Σ (a standard way of representing a set of TGDs as a graph, which is defined in Section 3) does not contain a “bad” cycle, where a “bad” cycle witnesses the fact that during the chase of D with Σ we eventually fall in a cyclic chase derivation that leads to non-termination. This again leads to conceptually simple chase termination algorithms: construct the dependency graph of Σ , and if it has a “bad” cycle, then conclude that the chase does not terminate; otherwise, it does.

We performed an exploratory analysis revealing that the algorithms based on materialization are too expensive in practice; note that our implementation of the materialization-based algorithms exploits VLog, a state-of-the-art chase engine [19]. We considered 483 pairs (D, Σ) , where D is a database and Σ a set of simple-linear TGDs, obtained from the ontology repository of the Data and Knowledge Group of the University of Oxford by keeping from the available ontologies, which are actually modelled using Description Logics, only the ontological axioms that can be expressed as simple-linear TGDs. For each such pair, we checked whether the semi-oblivious chase of D with Σ terminates by using both the materialization- and the acyclicity-based algorithms. We observed that for 28% of the scenarios the materialization-based algorithm runs out of memory, whereas the acyclicity-based algorithm provides an answer in 0.5 seconds in the worst-case. This is because the worst-case upper bound on the size of the result of the chase from [5] is very large, and thus, the materialization-based algorithm is forced, in general,

to construct extremely large instances before being able to safely recognize that the chase does not terminate. For the remaining 72% of the scenarios, the acyclicity-based algorithm consistently outperforms the materialization-based algorithm. Overall, for 50% of the considered scenarios, either the materialization-based algorithm fails, or the acyclicity-based algorithm is at least 10 times faster. Therefore, we concentrate on the acyclicity-based algorithms.

Main Outcome and Challenges. Our experimental analysis revealed that for simple-linear TGDs the primary parameter impacting the runtime of the acyclicity-based algorithm is the size of the input set of TGDs, whereas the size of the input database does not play any crucial role. Interestingly, the algorithm is very fast (in the order of seconds) even for large sets of simple-linear TGDs (with up to 200K TGDs). Now, concerning the more interesting case of linear TGDs, our analysis showed that the acyclicity-based algorithm consists of two components that are of different nature. In particular, there is a database-dependent component, whose performance is solely impacted by the size of the database, and a database-independent component, whose runtime is primarily affected by the size of the set of TGDs. Interestingly, the overall runtime of the algorithm is quite reasonable, which is a strong evidence that fast checking for the termination of the semi-oblivious chase in the presence of linear TGDs is not an unrealistic goal. Note that most of the total end-to-end runtime of the algorithm is taken by the database-dependent component, which indicates that our future efforts should be focused on improving that component.

Towards the above outcome concerning the acyclicity-based algorithms, we had to overcome a couple of technical challenges that led to results of independent interest:

- It would not be possible to obtain the above insightful conclusions by naively implementing the algorithms in question as this would lead to poor performance; this is further discussed in Section 4. Hence, we had to revisit and refine the theoretical algorithms from [5] in order to obtain novel algorithms that are amenable to efficient implementations. The low-level implementation details of the new refined algorithms are discussed in Section 5.

- In order to stress test the algorithms in question, we had to synthetically generate databases and sets of TGDs. However, as discussed in Section 6, existing data and TGD generators are not suitable for our purposes as they do not allow us to tune certain parameters that are crucial for evaluating our chase termination algorithms. To this end, we developed our own data and TGD generators, and used them to carefully generate the databases and TGDs that have been employed in our experimental analysis.

2 PRELIMINARIES

We consider the disjoint countably infinite sets $\mathbf{C}$, $\mathbf{N}$, and $\mathbf{V}$ of *constants*, (*labeled*) *nulls*, and *variables*, respectively. We refer to constants, nulls and variables as *terms*. For an integer $n > 0$, we write $[n]$ for the set of integers $\{1, \dots, n\}$.

Relational Databases. A *schema* $\mathbf{S}$ is a finite set of relation symbols (or predicates) with associated arity. We write R/n to denote that R has arity $n > 0$; we may also write $\text{ar}(R)$ for the integer n . A (*predicate*) *position* of $\mathbf{S}$ is a pair (R, i) , where $R/n \in \mathbf{S}$ and $i \in [n]$, that essentially identifies the i -th argument of R . We write $\text{pos}(\mathbf{S})$ for

the set of positions of $\mathbf{S}$, that is, the set $\{(R, i) \mid R/n \in \mathbf{S} \text{ and } i \in [n]\}$. An *atom* over $\mathbf{S}$ is an expression of the form $R(\bar{t})$, where $R/n \in \mathbf{S}$ and $\bar{t}$ is an n -tuple of terms. A *fact* is an atom whose arguments consist only of constants. For a variable x in $\bar{t} = (t_1, \dots, t_n)$, let $\text{pos}(R(\bar{t}), x) = \{(R, i) \mid t_i = x\}$. We write $\text{var}(R(\bar{t}))$ for the set of variables in $\bar{t}$. The notations $\text{pos}(\cdot, x)$ and $\text{var}(\cdot)$ extend to sets of atoms. An *instance* over $\mathbf{S}$ is a (possibly infinite) set of atoms over $\mathbf{S}$ with constants and nulls. A *database* over $\mathbf{S}$ is a finite set of facts over $\mathbf{S}$. The *active domain* of an instance I , denoted $\text{dom}(I)$, is the set of terms (constants and nulls) occurring in I .

Homomorphisms. A *homomorphism* from a set of atoms A to a set of atoms B is a function h from the set of terms in A to the set of terms in B such that h is the identity on $\mathbf{C}$, and $R(t_1, \dots, t_n) \in A$ implies $R(h(t_1), \dots, h(t_n)) \in B$.

Tuple-Generating Dependencies. A *tuple-generating dependency* (TGD) σ is a (constant-free) sentence $\forall \bar{x} \forall \bar{y} (\phi(\bar{x}, \bar{y}) \rightarrow \exists \bar{z} \psi(\bar{x}, \bar{z}))$, where $\bar{x}, \bar{y}$ and $\bar{z}$ are tuples of variables of $\mathbf{V}$, and $\phi(\bar{x}, \bar{y})$ and $\psi(\bar{x}, \bar{z})$ are non-empty conjunctions of atoms that mention only variables from $\bar{x} \cup \bar{y}$ and $\bar{x} \cup \bar{z}$, respectively. Note that, by abuse of notation, we may treat a tuple of variables as a set of variables. We write σ as $\phi(\bar{x}, \bar{y}) \rightarrow \exists \bar{z} \psi(\bar{x}, \bar{z})$, and use comma instead of $\wedge$ for joining atoms. We refer to $\phi(\bar{x}, \bar{y})$ and $\psi(\bar{x}, \bar{z})$ as the *body* and *head* of σ , denoted $\text{body}(\sigma)$ and $\text{head}(\sigma)$, respectively. The *frontier* of the TGD σ , denoted $\text{fr}(\sigma)$, is the set of variables $\bar{x}$, i.e., the variables that appear both in the body and the head of σ . The *schema* of a set Σ of TGDs, denoted $\text{sch}(\Sigma)$, is the set of predicates occurring in Σ . An instance I satisfies a TGD σ as the one above, written $I \models \sigma$, if whenever there exists a homomorphism h from $\phi(\bar{x}, \bar{y})$ to I , then there is an extension of h that is a homomorphism from $\psi(\bar{x}, \bar{z})$ to I ; we may treat a conjunction of atoms as a set of atoms. The instance I satisfies a set Σ of TGDs, written $I \models \Sigma$, if $I \models \sigma$ for each $\sigma \in \Sigma$.

Linearity. A TGD is called *linear* if it has only one body-atom, and the corresponding class that collects all the finite sets of linear TGDs is denoted $\mathbf{L}$. We call a linear TGD *simple* if no variable occurs more than once in its body, and the obtained class is denoted $\mathbf{SL}$.

3 THE SEMI-OBLIVIOUS CHASE PROCEDURE

The semi-oblivious chase (or simply chase) takes as input a database D and a set Σ of TGDs, and constructs an instance that contains D and satisfies Σ . A central notion in this context is that of trigger.

Definition 3.1. Given a set Σ of TGDs and an instance I , a *trigger* for Σ on I is a pair (σ, h) , where $\sigma \in \Sigma$ and h is a homomorphism from $\text{body}(\sigma)$ to I . The *result* of (σ, h) , denoted $\text{result}(\sigma, h)$, is the set $\mu(\text{head}(\sigma))$, where $\mu : \text{var}(\text{head}(\sigma)) \rightarrow \mathbf{C} \cup \mathbf{N}$ is defined as follows: $\mu(x) = h(x)$ if $x \in \text{fr}(\sigma)$, and $\mu(x) = \perp_{\sigma, h|_{\text{fr}(\sigma)}}^x$ if $x \notin \text{fr}(\sigma)$, where $\perp_{\sigma, h|_{\text{fr}(\sigma)}}^x$ is a null. Let $T(\Sigma, I)$ be the set of triggers for Σ on I . ■

Observe that in the definition of $\text{result}(\sigma, h)$, each existentially quantified variable x of $\text{head}(\sigma)$ is mapped by μ to a null value of $\mathbf{N}$ whose name is uniquely determined by the trigger (σ, h) and the variable x itself. This means that, given a trigger (σ, h) , we can unambiguously construct the set of atoms $\text{result}(\sigma, h)$. The central idea of the chase is, starting from a database D , to exhaustively apply triggers for the given set Σ of TGDs on the instance constructed

so far. More precisely, given a database D and a set Σ of TGDs, let $\text{chase}^0(D, \Sigma) = D$, and for each $i > 0$, let

$$\text{chase}^i(D, \Sigma) = \text{chase}^{i-1}(D, \Sigma) \cup \bigcup_{(\sigma, h) \in S} \text{result}(\sigma, h),$$

where $S = T(\Sigma, \text{chase}^{i-1}(D, \Sigma))$. We finally define *the result of the chase of D w.r.t. Σ* as the instance $\text{chase}(D, \Sigma) = \bigcup_{i \geq 0} \text{chase}^i(D, \Sigma)$.

Chase Termination. The result of the chase may be infinite even for very simple settings: it is easy to see that for $D = \{R(a, b)\}$ and $\Sigma = \{R(x, y) \rightarrow \exists z R(y, z)\}$, $\text{chase}(D, \Sigma)$ is infinite. This leads to the following problem, parameterized by a class C of TGDs:

INPUT : A database D and a set Σ of TGDs from C .
 QUESTION : Is the instance $\text{chase}(D, \Sigma)$ finite?

This problem has been recently studied in [5] for the classes of simple-linear and linear TGDs. Interestingly, for both classes, the finiteness of the result of the chase has been syntactically characterized by exploiting the notion of non-uniform weak-acyclicity. We proceed to recall this acyclicity notion, and then present the characterizations established in [5], which in turn lead to simple algorithms for checking the finiteness of the chase. Note that, for clarity, in the rest of the paper we assume TGDs with a non-empty frontier, i.e., we assume that in each TGD σ there is at least one variable that occurs both in $\text{body}(\sigma)$ and $\text{head}(\sigma)$. This assumption can be made without loss of generality since, given a database D and a set Σ of TGDs, we can easily construct a set Σ' of TGDs with a non-empty frontier by slightly modifying Σ such that $\text{chase}(D, \Sigma)$ is finite iff $\text{chase}(D, \Sigma')$ is finite.

Non-Uniform Weak-Acyclicity. Weak-acyclicity was introduced in [10] as the main formalism for data exchange purposes, which guarantees the finiteness of the result of the chase for *every* input database. Non-uniform weak-acyclicity is the database-dependent variant of weak-acyclicity introduced in [5]. We proceed to give the formal definitions. We first need to recall the notion of the *dependency graph* of a set Σ of TGDs, defined as a directed multigraph $\text{dg}(\Sigma) = (N, E)$, where $N = \text{pos}(\text{sch}(\Sigma))$ and E contains *only* the following edges. For each TGD $\sigma \in \Sigma$ with $\text{head}(\sigma) = \{\alpha_1, \dots, \alpha_k\}$, for each $x \in \text{fr}(\sigma)$, and for each position $\pi \in \text{pos}(\text{body}(\sigma), x)$:

- For each $i \in [k]$ and for each $\pi' \in \text{pos}(\alpha_i, x)$, there exists a *normal* edge $(\pi, \pi') \in E$.
- For each existentially quantified variable z in σ , $i \in [k]$, and $\pi' \in \text{pos}(\alpha_i, z)$, there is a *special* edge $(\pi, \pi') \in E$.

We further need to define when a predicate is reachable from another predicate. Given predicates $R, P \in \text{sch}(\Sigma)$, P is *reachable from R* (w.r.t. Σ) if $R = P$, or there exists a path in $\text{dg}(\Sigma)$ from a position of the form (R, i) to a position of the form (P, j) . Given a database D , we say that a (not necessarily simple and possibly cyclic) path C in $\text{dg}(\Sigma)$ is *D-supported* if there exists an atom $R(\bar{t}) \in D$ and a node of the form (P, i) in C such that P is reachable from R . We are now ready to recall (non-uniform) weak-acyclicity.

Definition 3.2. Consider a database D and a set Σ of TGDs. We say that Σ is *weakly-acyclic w.r.t. D* , or *D-weakly-acyclic*, if there is no D -supported cycle in $\text{dg}(\Sigma)$ with a special edge. We say that Σ is *weakly-acyclic* if there is no cycle in $\text{dg}(\Sigma)$ with a special edge. ■

Characterizing the Finiteness of the Chase. It is not very difficult to show that whenever a set Σ of TGDs (not necessarily linear) is D -weakly-acyclic, then the instance $\text{chase}(D, \Sigma)$ is finite. In other words, the D -weak-acyclicity of Σ is a sufficient condition for the finiteness of $\text{chase}(D, \Sigma)$. What is more interesting is that, assuming that Σ is a set of simple-linear TGDs, the D -weak-acyclicity of Σ is also a necessary condition for the finiteness of $\text{chase}(D, \Sigma)$. This leads to the following characterization established in [5]:

THEOREM 3.3. Consider a database D and a set $\Sigma \in \text{SL}$ of TGDs. It holds that $\text{chase}(D, \Sigma)$ is finite iff Σ is D -weakly-acyclic.

As shown in [5], non-uniform weak-acyclicity is not powerful enough for characterizing the finiteness of the chase instance in the case of linear TGDs. To obtain a characterization analogous to Theorem 3.3, the authors of [5] used the technique of *simplification* to convert linear TGDs into simple-linear TGDs, while preserving the finiteness of the chase instance. We recall this technique. Let $\bar{t} = (t_1, \dots, t_n)$ be a tuple of (not necessarily distinct) terms. We write $\text{unique}(\bar{t})$ for the tuple obtained from $\bar{t}$ by keeping only the first occurrence of each term in $\bar{t}$. For example, if $\bar{t} = (x, y, x, z, y)$, then $\text{unique}(\bar{t}) = (x, y, z)$. For each $i \in [n]$, the *identifier of t_i in $\bar{t}$* , denoted $\text{id}_{\bar{t}}(t_i)$, is the integer that identifies the position of $\text{unique}(\bar{t})$ at which t_i appears. We write $\text{id}(\bar{t})$ for the tuple $(\text{id}_{\bar{t}}(t_1), \dots, \text{id}_{\bar{t}}(t_n))$. For example, if $\bar{t} = (x, y, x, z, y)$, then $\text{id}(\bar{t}) = (1, 2, 1, 3, 2)$. For an atom $\alpha = R(\bar{t})$, the *shape of α* , denoted $\text{shape}(\alpha)$, is the predicate $R_{\text{id}(\bar{t})}$, whereas the *simplification of α* , denoted $\text{simple}(\alpha)$, is the atom $R_{\text{id}(\bar{t})}(\text{unique}(\bar{t}))$. For example, assuming that α is $R(x, y, x, z)$, $\text{shape}(\alpha)$ is the predicate $R_{(1,2,1,3)}$, which essentially describes how the terms occurring in α are related, that is, there are three distinct terms, while the first and the third are the same. The simplification of α is the atom $R_{(1,2,1,3)}(x, y, z)$, which fully describes the atom α without repeating variables. Indeed, given $R_{(1,2,1,3)}(x, y, z)$, we can unambiguously construct the atom of which $R_{(1,2,1,3)}(x, y, z)$ is the simplification of, that is, the atom with predicate R having at its first and third position the first variable (i.e., x), at its second position the second variable (i.e., y), and at its fourth position the third variable (i.e., z). We can naturally refer to the simplification and the shape of a set of atoms. For a tuple of variables $\bar{x} = (x_1, \dots, x_n)$, a *specialization of $\bar{x}$* is a function f from $\bar{x}$ to $\bar{x}$ such that $f(x_1) = x_1$, and $f(x_i) \in \{f(x_1), \dots, f(x_{i-1}), x_i\}$, for each $i \in \{2, \dots, n\}$. We write $f(\bar{x})$ for $(f(x_1), \dots, f(x_n))$. We are now ready to recall how a set of linear TGDs is converted into a set of simple-linear TGDs.

Definition 3.4. Consider a linear TGD σ of the form $R(\bar{x}) \rightarrow \exists \bar{z} \psi(\bar{y}, \bar{z})$, where $\bar{y} \subseteq \bar{x}$, and a specialization f of $\bar{x}$. The *simplification of σ induced by f* is the simple-linear TGD

$$\text{simple}(R(f(\bar{x}))) \rightarrow \exists \bar{z} \text{simple}(\psi(f(\bar{y}), \bar{z})).$$

We write $\text{simple}(\sigma)$ for the set of all simplifications of σ induced by some specialization of $\bar{x}$. For a set $\Sigma \in \text{L}$ of TGDs, the *simplification of Σ* is defined as the set $\text{simple}(\Sigma) = \bigcup_{\sigma \in \Sigma} \text{simple}(\sigma)$. ■

We proceed to give a simple example that illustrates the notions of specialization and simplification introduced above.

Example 3.5. Consider the set Σ consisting of the linear TGDs

$$\sigma_1 : R(x, y, x, z) \rightarrow \exists w Q(x, w) \quad \sigma_2 : Q(x, y) \rightarrow R(x, x, y, y).$$

For a database D , during the construction of $\text{chase}(D, \Sigma)$, the TGD σ_1 can be triggered due to an atom α of the form $R(t_1, t_2, t_3, t_4)$ as long as $t_1 = t_3$. Hence, even if $t_1 = t_2 = t_3$ or $t_1 = t_2 = t_3 = t_4$, the atom α also triggers σ_1 . The goal of simplifying a linear TGD, for example σ_1 , is to convert it into a set of linear TGDs that account for all possible ways of triggering σ_1 , but without repeating variables in their bodies, i.e., they are simple-linear. Each specialization f of the body-variables of σ_1 encodes one way of triggering σ_1 . For example, the function f such that $f(x) = f(y) = x$ and $f(z) = z$ induces the atom $R(x, x, x, z)$, which represents a “more specific” way of triggering σ_1 w.r.t. the original body-atom $R(x, y, x, z)$, since it not only requires the first and third position of the atom to unify to the same term, but the second position as well. Similarly, if f is such that $f(x) = f(y) = f(z) = x$, then it encodes the atom $R(x, x, x, x)$ that is even more specific. Hence, the simplification $\text{simple}(\sigma_1)$ of σ_1 consists of the following simple-linear TGDs:

$$\begin{aligned} R_{(1,2,1,3)}(x, y, z) &\rightarrow \exists w Q_{(1,2)}(x, w) \\ R_{(1,1,1,2)}(x, z) &\rightarrow \exists w Q_{(1,2)}(x, w) \\ R_{(1,2,1,1)}(x, y) &\rightarrow \exists w Q_{(1,2)}(x, w) \\ R_{(1,2,1,2)}(x, y) &\rightarrow \exists w Q_{(1,2)}(x, w) \\ R_{(1,1,1,1)}(x) &\rightarrow \exists w Q_{(1,2)}(x, w). \end{aligned}$$

The above simple-linear TGDs take into account all possible ways σ_1 can be triggered during the construction of the chase; the head atoms are simplified accordingly. Similarly, $\text{simple}(\sigma_2)$, the simplification of σ_2 , consists of the following simple-linear TGDs:

$$Q_{(1,2)}(x, y) \rightarrow R_{(1,1,2,2)}(x, y) \quad Q_{(1,1)}(x) \rightarrow R_{(1,1,1,1)}(x).$$

The chase of D w.r.t. Σ can be faithfully simulated via the chase of $\text{simple}(D)$ w.r.t. $\text{simple}(\Sigma) = \text{simple}(\sigma_1) \cup \text{simple}(\sigma_2)$. ■

We can now recall the characterization for the finiteness of the chase instance for linear TGDs, established in [5], which is similar to the one for simple-linear TGDs, with the key difference that first we need to simplify both the database and the set of linear TGDs:

THEOREM 3.6. *Consider a database D and a set $\Sigma \in \mathcal{L}$ of TGDs. Then, $\text{chase}(D, \Sigma)$ is finite iff $\text{simple}(\Sigma)$ is $\text{simple}(D)$ -weakly-acyclic.*

It is clear that Theorems 3.3 and 3.6 provide simple algorithms for checking whether the chase instance is finite. Our goal is to experimentally evaluate those algorithms with the aim of understanding which input parameters affect their performance, clarifying whether they can be applied in a practical context, and revealing their performance limitations. Of course, a naive implementation of the obtained algorithms, especially for linear TGDs where the expensive simplification must be applied, will lead to poor performance, and thus, will not be very useful towards our goal. Indeed, simplification is expensive as, in general, it leads to a very large number of simple-linear TGDs. Consider, for example, the TGD σ of the form $P(x, x, y_1, \dots, y_n) \rightarrow T(x)$, for some $n > 0$. One can verify that the number of TGDs in $\text{simple}(\sigma)$ coincides with the number of ways one can partition a set of n elements. This is known as the n -th Bell number, denoted B_n , where $B_0 = 1$ and $B_{i+1} = \sum_{k=0}^i \binom{i}{k} B_k$. Thus, $B_n = |\text{simple}(\sigma)| \geq 2^{n-1}$, i.e., is at least exponential in n . Hence, we need to convert the inherited theoretical algorithms into practical algorithms that are amenable to efficient implementations.

Algorithm 1: IsChaseFinite[SL]

Input: A database D and a set $\Sigma \in \mathcal{SL}$ of TGDs

Output: true if $\text{chase}(D, \Sigma)$ is finite and false otherwise

```

1  $G \leftarrow \text{BuildDepGraph}(\Sigma);$ 
2  $S \leftarrow \text{FindSpecialSCC}(G);$ 
3  $P \leftarrow \bigcup_{C \in S} \{v_C\};$ 
4 if Supports( $D, P, G$ ) then return false;
5 return true

```

4 PRACTICAL TERMINATION ALGORITHMS

We first present the algorithm IsChaseFinite[SL] that accepts as input a database D and a set Σ of simple-linear TGDs, and checks whether Σ is D -weakly-acyclic, i.e., whether $\text{chase}(D, \Sigma)$ is finite. Note that a naive search for a “bad” cycle in a dependency graph will be too costly since we may have to go through exponentially many cycles. Thus, IsChaseFinite[SL] relies on a refined machinery that searches for *strongly connected components* with a special edge. We then proceed to give an analogous algorithm, dubbed IsChaseFinite[L] , for linear TGDs, which essentially simplifies the given database D and set Σ of linear TGDs, and then checks whether $\text{simple}(\Sigma)$ is $\text{simple}(D)$ -weakly-acyclic, which is equivalent to say that $\text{chase}(D, \Sigma)$ is finite. Note, however, that IsChaseFinite[L] relies on a refined notion of simplification that *dynamically simplifies* Σ by leveraging the given database D , instead of doing it statically as in Definition 3.4. For instance, consider the set Σ of TGDs from Example 3.5 and the database $D = \{R(a, b, a, c)\}$. We have that $\text{simple}(D) = \{R_{(1,2,1,3)}(a, b, c)\}$. It is clear that most of the TGDs of $\text{simple}(\Sigma)$ are not needed to construct $\text{chase}(\text{simple}(D), \text{simple}(\Sigma))$ as the only TGDs being triggered, starting from $\text{simple}(D)$, are $R_{(1,2,1,3)}(x, y) \rightarrow \exists w Q_{(1,2)}(x, w)$ and $Q_{(1,2)}(x, y) \rightarrow R_{(1,1,2,2)}(x, y)$. The goal of dynamic simplification is to keep only TGDs from $\text{simple}(\Sigma)$ that are needed to construct $\text{chase}(\text{simple}(D), \text{simple}(\Sigma))$. We present the above algorithms at a high-level, whereas their implementation details are discussed in Section 5.

4.1 Simple-Linear TGDs

A *strongly connected component* (SCC) in a directed graph G is a maximal subgraph of G in which there is a (directed) path between every pair of nodes. A *special SCC* in a dependency graph is an SCC with at least one special edge. IsChaseFinite[SL] , which is depicted in Algorithm 1, starts by building the dependency graph G of the input set Σ of TGDs (line 1). It then collects the special SCCs of G in a set S (line 2), which can clearly form “bad” cycles that violate non-uniform weak-acyclicity. Of course, for the latter to happen, some nodes (i.e., predicate positions) in a special SCC must be supported by the given database D as defined in Section 2. To check this, the algorithm first collects exactly one node v_C from each special SCC C of G in a set P (line 3); it is not important how v_C is selected. It then checks if D supports any of the nodes of P (line 4). If this is the case, then there is a D -supported cycle in G with a special edge, and thus the algorithm returns false; otherwise, it returns true. The correctness of IsChaseFinite[SL] follows by Theorem 3.3:

LEMMA 4.1. *Consider a database D and a set $\Sigma \in \mathcal{SL}$ of TGDs. It holds that $\text{IsChaseFinite[SL]}(D, \Sigma) = \text{true}$ iff $\text{chase}(D, \Sigma)$ is finite.*

4.2 Linear TGDs

Although the algorithm `IsChaseFinite[SL]` together with the simplification technique (see Definition 3.4) immediately give rise to a simple algorithm for checking the finiteness of the chase instance for linear TGDs, a naive implementation of the simplification technique leads to poor performance. Indeed, we performed exploratory experiments on sets of linear TGDs coming from the literature and observed that a naive implementation is not scalable as the algorithm quickly runs out of memory when dealing with large sets of TGDs. This is because by statically simplifying a set of linear TGDs Σ , without taking into account the underlying database, leads to an exponentially large set of simple-linear TGDs; in particular, the size of the set $\text{simple}(\Sigma)$ is exponential in the maximum arity of the predicates in $\text{sch}(\Sigma)$. Thus, the algorithm `IsChaseFinite[SL]` becomes impractical due to the very large size of the dependency graph of $\text{simple}(\Sigma)$, which exceeds the capacity of the main memory.

Dynamic Simplification. We refine the notion of simplification by taking into account the underlying database, which leads to the technique of dynamic simplification. In particular, given a database D and a set Σ of linear TGDs, the goal is to define a set $\text{simple}_D(\Sigma)$, which is a subset of $\text{simple}(\Sigma)$, that enjoys two crucial properties:

- (1) It holds that the instance $\text{chase}(\text{simple}(D), \text{simple}(\Sigma))$ is finite iff the instance $\text{chase}(\text{simple}(D), \text{simple}_D(\Sigma))$ is finite, which essentially tells us that the technique of dynamic simplification preserves the finiteness of the chase.
- (2) The set $\text{simple}_D(\Sigma)$ is, generally, orders of magnitude smaller than the set $\text{simple}(\Sigma)$ obtained by statically simplifying Σ .

Item (1) is established by Lemma 4.4 below. Item (2) cannot be mathematically proved as there are cases where both static and dynamic simplification build the same set of linear TGDs. However, we have experimentally verified that for existing databases and sets of TGDs coming from the literature (in fact, those used in Section 9), the size of the dynamically simplified sets of TGDs is, on average, 5 times smaller than the size of the corresponding statically simplified sets of TGDs. The absolute difference varies with the dynamically simplified sets being up to 1000 times smaller in the best case.

The key idea of dynamic simplification is to exploit the shapes of the atoms occurring in the given database to guide the simplification. More precisely, given a database D and a set Σ of linear TGDs, we first collect the shapes that can be derived from $\text{shape}(D)$ using the TGDs of Σ ; we denote this set as $\Sigma(\text{shape}(D))$. Then, $\text{simple}_D(\Sigma)$ keeps from the set $\text{simple}(\Sigma)$ only those simple-linear TGDs such that the predicate of their body-atom belongs to $\Sigma(\text{shape}(D))$, as these are the only TGDs that can be applied during the construction of the instance $\text{chase}(\text{simple}(D), \text{simple}(\Sigma))$. All the other TGDs of $\text{simple}(\Sigma)$ are superfluous whenever the input database is D in the sense that they will never be applied during the construction of $\text{chase}(\text{simple}(D), \text{simple}(\Sigma))$. We proceed to formalize this idea. To this end, we need to introduce some auxiliary notions.

For a schema S , we use $\text{shape}(S)$ to denote the set of all shapes mentioning a predicate of S . For a set of shapes $S \subseteq \text{shape}(S)$, the *database induced by S* , denoted $DB[S]$, is the database of the form $\{R(\text{id}(\bar{i})) \mid R_{\text{id}(\bar{i})} \in S\}$. For example, assuming that $S = \{R_{(1,2)}, P_{(1,1,2)}\}$, then $DB[S] = \{R(1, 2), P(1, 1, 2)\}$. Consider now a linear TGD $\sigma = R(x_1, \dots, x_n) \rightarrow \exists \bar{z} \psi(\bar{y}, \bar{z})$ and let h be a homomorphism from $\{R(x_1, \dots, x_n)\}$ to $\{R(i_1, \dots, i_n)\} \subseteq DB[\text{shape}(\{R\})]$.

The *h -specialization* of the tuple $(x_1, \dots, x_n)$ is the (unique) specialization f of $(x_1, \dots, x_n)$ such that $f(x_i) = f(x_j)$ iff $h(x_i) = h(x_j)$, for every $i, j \in [n]$. For example, if h is a homomorphism from $\{R(x, y, x, z)\}$ to $\{R(1, 1, 1, 2)\}$, the h -specialization of (x, y, x, z) is the function f such that $f(x) = x$, $f(y) = x$, and $f(z) = z$. We can now proceed with the formalization of dynamic simplification.

Consider a set Σ of linear TGDs and a set of shapes $S \subseteq \text{shape}(\Sigma)$; for brevity, we write $\text{shape}(\Sigma)$ for $\text{shape}(\text{sch}(\Sigma))$. A shape $R_{\text{id}(\bar{i})} \in \text{shape}(\Sigma)$ is an *immediate consequence* of S and Σ if:

- (1) $R_{\text{id}(\bar{i})} \in S$, or
- (2) there is a TGD $P(\bar{x}) \rightarrow \exists \bar{z} \psi(\bar{y}, \bar{z})$ in Σ and a homomorphism h from $\{P(\bar{x})\}$ to $DB[S]$ such that $R_{\text{id}(\bar{i})}$ occurs in the head of the simplification of σ induced by the h -specialization of $\bar{x}$.

In simple words, item (2) tells us that there exists a TGD in $\text{simple}(\Sigma)$ of the form $P_{\text{id}(\bar{i}')}(\bar{x}) \rightarrow \exists \bar{z} \dots, R_{\text{id}(\bar{i})}(\bar{y}), \dots$ with $P_{\text{id}(\bar{i}')} \in S$. The *immediate consequence operator* of Σ is the function $\Gamma_\Sigma : 2^{\text{shape}(\Sigma)} \rightarrow 2^{\text{shape}(\Sigma)}$ (as usual, 2^X denotes the powerset of a set X) such that

$$\Gamma_\Sigma(S) = \{R_{\text{id}(\bar{i})} \mid R_{\text{id}(\bar{i})} \text{ is an immediate consequence of } S \text{ and } \Sigma\}.$$

By iterative applications of the above operator, we can compute the shapes that can be derived from S using the TGDs of Σ . Formally, $\Gamma_\Sigma^0(S) = S$ and $\Gamma_\Sigma^i(S) = \Gamma_\Sigma(\Gamma_\Sigma^{i-1}(S))$, for each $i > 0$, and we finally let $\Sigma(S) = \bigcup_{i \geq 0} \Gamma_\Sigma^i(S)$. At first glance, the construction of $\Sigma(S)$ requires infinitely many iterations. However, since $\Sigma(S) \subseteq \text{shape}(\Sigma)$, in the worst-case $\Sigma(S)$ is obtained after $|\text{shape}(\Sigma)|$ iterations. It is actually easy to verify that $\Sigma(S) = \Gamma_\Sigma^{|\text{simple}(\Sigma)|}(S)$. Therefore, since $\text{shape}(\Sigma)$ is finite, $\Sigma(S)$ can be obtained after finitely many steps. We now have all the ingredients to formally define dynamic simplification.

Definition 4.2. Consider a database D and a set Σ of linear TGDs.¹ The *dynamic simplification of Σ relative to D* (or *D -simplification of Σ*), denoted $\text{simple}_D(\Sigma)$, is defined as the set

$$\{\text{simple}(R(f(\bar{x}))) \rightarrow \exists \bar{z} \text{simple}(\psi(f(\bar{y}), \bar{z})) \mid \\ R(\bar{x}) \rightarrow \exists \bar{z} \psi(\bar{y}, \bar{z}) \in \Sigma \text{ and } f \text{ is the } h\text{-specialization of } \bar{x} \\ \text{for some homomorphism } h \text{ from } \{R(\bar{x})\} \text{ to } DB(\Sigma(\text{shape}(D)))\}$$

consisting only of simple-linear TGDs. ■

It is not difficult to verify that the D -simplification of Σ essentially collects all the TGDs of $\text{simple}(\Sigma)$ such that the predicate of their body-atom belongs to $\Sigma(\text{shape}(D))$. A simple example that illustrates the notion of dynamic simplification follows.

Example 4.3. Consider again the set Σ of TGDs from Example 3.5 and the database $D = \{R(a, b, a, c)\}$. We proceed to describe how the dynamic simplification procedure operates on D and Σ . The sets S and ΔS are initialized to the set of shapes of D , that is, $S = \Delta S = \{R_{(1,2,1,3)}\}$, and $\Sigma_s = \emptyset$. At the first iteration of the algorithm we have that σ_1 is “applicable” to ΔS , i.e., the body-atom of σ_1 is “coherent” with the (only) shape present in ΔS , i.e., $R_{(1,2,1,3)}$. No other TGDs of Σ are applicable, and thus, we conclude that the TGD $R_{(1,2,1,3)}(x, y, z) \rightarrow \exists w Q_{(1,2)}(x, w)$, which is constructed in line 5 and added to Σ_s in line 7, will be actually triggered by the chase. We then have $S = \{R_{(1,2,1,3)}, Q_{(1,2)}\}$, while the new

¹We assume, without loss of generality, that the atoms of D mention only predicates of $\text{sch}(\Sigma)$, and thus, $\text{shape}(D) \subseteq \text{shape}(\Sigma)$. Indeed, the atoms of D with a predicate not in $\text{sch}(\Sigma)$ do not affect in any way the size of the instance $\text{chase}(D, \Sigma)$.

Algorithm 2: DynSimplification

Input: A database D and a set $\Sigma \in \mathbf{L}$ of TGDs

Output: The D -simplification of Σ

```
1  $S \leftarrow \text{FindShapes}(D)$ ;  
2  $\Sigma_S \leftarrow \emptyset$ ;  
3  $\Delta S \leftarrow S$ ;  
4 while  $\Delta S \neq \emptyset$  do  
5    $\Sigma_{aux} \leftarrow \text{Applicable}(\Delta S, \Sigma)$ ;  
6    $S_{aux} \leftarrow \{R_{id(\bar{i})} \in \text{shape}(\Sigma) \mid \text{there exists a TGD } \sigma \in \Sigma_{aux} \text{ such that } R_{id(\bar{i})} \text{ occurs in head}(\sigma)\}$ ;  
7    $\Sigma_S \leftarrow \Sigma_S \cup \Sigma_{aux}$ ;  
8    $\Delta S \leftarrow S_{aux} \setminus S$ ;  
9    $S \leftarrow S \cup \Delta S$ ;  
10 return  $\Sigma_S$ ;
```

shapes are $\Delta S = \{Q_{(1,2)}\}$. At the second iteration we have that σ_2 is applicable to ΔS . Thus, the TGD $Q_{(1,2)}(x, y) \rightarrow R_{(1,1,2,2)}(x, y)$ is added to Σ_S and we have $S = \{R_{(1,2,1,3)}, Q_{(1,2)}, R_{(1,1,2,2)}\}$ and $\Delta S = \{R_{(1,1,2,2)}\}$. Finally, at the third iteration, no TGDs of Σ are applicable since there is no way to specialize the body-atom of σ_1 in a way that it becomes coherent with the shape $R_{(1,1,2,2)}$ in ΔS . The algorithm then terminates and Σ_S collects a subset of $\text{simple}(\Sigma)$ that is sufficient to construct $\text{chase}(\text{simple}(D), \text{simple}(\Sigma))$. ■

We now proceed to show that indeed dynamic simplification preserves the finiteness of the chase.

LEMMA 4.4. *Consider a database D and a set $\Sigma \in \mathbf{L}$ of TGDs. The following are equivalent:*

- (1) $\text{chase}(\text{simple}(D), \text{simple}(\Sigma))$ is finite.
- (2) $\text{chase}(\text{simple}(D), \text{simple}_D(\Sigma))$ is finite.

Since, by definition, $\text{simple}_D(\Sigma) \subseteq \text{simple}(\Sigma)$, it is clear that (1) implies (2) holds trivially. The interesting direction is (2) implies (1), which can be shown via an inductive argument. Another crucial property of dynamic simplification, which will help us to further improve the performance of the termination algorithm, is that the $\text{simple}(D)$ -weak-acyclicity of $\text{simple}_D(\Sigma)$ coincides with the weak-acyclicity of $\text{simple}_D(\Sigma)$, which in turn leads to the following result:

LEMMA 4.5. *Consider a database D and a set $\Sigma \in \mathbf{L}$ of TGDs. The following are equivalent:*

- (1) $\text{chase}(\text{simple}(D), \text{simple}_D(\Sigma))$ is finite.
- (2) $\text{simple}_D(\Sigma)$ is weakly-acyclic.

An Algorithm for Dynamic Simplification. We now provide a concrete algorithm that performs the dynamic simplification of a set of linear TGDs that is amenable to an efficient implementation. To this end, we present the algorithm DynSimplification, depicted in Algorithm 2. The algorithm starts by finding the shapes of the atoms occurring in D , namely it computes the set $\text{shape}(D)$ (line 1). It then initializes the set of simplified TGDs Σ_S (line 2) and the set of new shapes ΔS (line 3). Then, the algorithm iteratively generates simplified TGDs and collects the new shapes that are added to ΔS , and continues this until a fixpoint is reached, i.e., $\Delta S = \emptyset$ (line 4). In particular, at each iteration, the algorithm computes simplified

Algorithm 3: IsChaseFinite[L]

Input: A database D and a set $\Sigma \in \mathbf{L}$ of TGDs

Output: true if $\text{chase}(D, \Sigma)$ is finite and false otherwise

```
1  $\Sigma_S \leftarrow \text{DynSimplification}(D, \Sigma)$ ;  
2  $G \leftarrow \text{BuildDepGraph}(\Sigma_S)$ ;  
3 if  $\text{FindSpecialSCC}(G) \neq \emptyset$  then return false;  
4 return true
```

TGDs that are not superfluous, i.e., they can be applied during the construction of $\text{chase}(\text{simple}(D), \text{simple}(\Sigma))$, that are added to Σ_S (lines 5 and 7). This is done via Applicable, which takes as input a set of shapes $\hat{S}$ and a set of linear TGDs $\hat{\Sigma}$, and returns the set

$$\{\text{simple}(R(f(\bar{x}))) \rightarrow \exists \bar{z} \text{simple}(\psi(f(\bar{y}), \bar{z})) \mid$$

$$R(\bar{x}) \rightarrow \exists \bar{z} \psi(\bar{y}, \bar{z}) \in \hat{\Sigma} \text{ and } f \text{ is the } h\text{-specialization of } \bar{x}$$

$$\text{for some homomorphism } h \text{ from } \{R(\bar{x})\} \text{ to } DB[\hat{S}]\}.$$

In essence, the procedure Applicable computes the set of TGDs of $\text{simple}(\hat{\Sigma})$ such that the predicate of their body belongs to $\hat{S}$. The algorithm also collects the newly generated shapes, that is, the predicates occurring in the head of the TGDs of $\text{Applicable}(\Delta S, \Sigma)$, that are added to ΔS (lines 6 and 8). Note that at each iteration, the algorithm applies the TGDs on ΔS , not on S , with the exception of the first iteration where $S = \Delta S$. This works because there are no new applicable TGDs on S after the first iteration since the TGDs are linear and all the applicable TGDs on S are applied during the first iteration. DynSimplification is correct by construction:

LEMMA 4.6. *Consider a database D and a set $\Sigma \in \mathbf{L}$ of TGDs. It holds that $\text{DynSimplification}(D, \Sigma) = \text{simple}_D(\Sigma)$.*

Termination Algorithm. Having in place DynSimplification, it is now straightforward to devise the algorithm IsChaseFinite[L], depicted in Algorithm 3, that checks for the finiteness of the chase in the case of linear TGDs. The correctness of DynSimplification, Theorem 3.6, Lemma 4.4, and Lemma 4.5, imply the correctness of the algorithm IsChaseFinite[L], and the next lemma follows:

LEMMA 4.7. *Given a database D and a set $\Sigma \in \mathbf{L}$ of TGDs, it holds that $\text{IsChaseFinite}[L](D, \Sigma) = \text{true}$ iff $\text{chase}(D, \Sigma)$ is finite.*

5 IMPLEMENTATION DETAILS

In this section, we give the essence of the key technical choices we made during the implementation of the algorithms from Section 4.

Build Dependency Graphs. The procedure BuildDepGraph takes as input a set Σ of TGDs, and returns the dependency graph of Σ in the form of an adjacency list, i.e., a list of lists, where each list corresponds to a node v , and the members in such a list represent the outgoing edges of v . Such an adjacency list is actually implemented as a doubly linked list. To speed up the process of building dependency graphs, the procedure also uses an index structure that maps predicate positions to their corresponding elements in the adjacency list. This index allows for fast access to the elements (nodes) for adding new links (edges) while parsing new TGDs.

Find Special SCCs. To implement FindSpecialSCC we adapt Tarjan's algorithm for finding SCCs in directed graphs [18]. Actually,

FindSpecialSCC is a simple extension of Tarjan’s algorithm. Such an extension is needed as we need a mechanism that allows us to check whether a SCC obtained by Tarjan’s algorithm is special.

Check for Positions Support. The procedure Supports consists of two steps: (1) query the database to find the positions of the predicates in the database, and (2) traverse the dependency graph starting from the positions in the special SCCs in the reverse order to reach the positions computed in the first step. Step (1) has been implemented via a single SQL query that returns the list of non-empty relations, which we then use to create the set of positions of the predicates occurring in the database, denoted P_D . For step (2), we start from the given set of positions P and traverse the graph in the reverse order using the reverse links in the adjacency list of the dependency graph. The procedure returns true if the graph traversal in the second step reaches a node (i.e., a position) in P_D ; otherwise, it returns false.

Dynamic Simplification. DynSimplification is an iterative procedure that uses FindShapes that computes the set of shapes S of the atoms of D and Applicable that computes the simplified TGDs using the shapes of S . We have two kinds of implementations for the procedure FindShapes, that is, *in-memory* and *in-database*. For the in-memory implementation, we run an SQL query for each relation R of D to load all the tuples of R into the main memory and then compute the shapes, whereas the in-database implementation does not load the relations, but instead runs SQL queries to find the shapes of the atoms in each relation. Concerning the procedure Applicable, recall that it takes as input a set S of shapes and a set Σ of linear TGDs, and returns the set of TGDs of simple(Σ) such that the predicate of their body belongs to S . The iterative process explained in Section 4 can be very costly with a large set of shapes and a large set of TGDs. To improve the performance, the implementation uses an index structure that enables fast access to the TGDs. Furthermore, for checking whether a relevant TGD is applicable, we generate and store the shape of the body-atom of each TGD. We also store an array of strings representing all possible identifiers of tuples up to the maximum arity of the schema that allows us to quickly find the shapes of the body-atoms in simplified TGDs.

6 EXPERIMENTAL INFRASTRUCTURE

Towards our goal, we are going to conduct extensive experiments with synthetic data and sets of TGDs. Therefore, we need a way to generate databases and sets of TGDs that are suitable for such an experimental evaluation. Note that existing data generators such as TPC-H and DataFiller, as well as TGD generators such as those from [2, 4], are not suitable for our purposes since they do not allow us to control the shape of atoms, which is crucial for generating databases and sets of TGDs that are suitable for our experiments. Thus, we had to implement our own data and TGD generators that support this key feature. The data generator takes as input a tuple of integer values for the tuning parameters ($preds$, min , max , $dsize$, $rsizes$), and randomly constructs a database D such that, with $S = \{R \mid R(\vec{c}) \in D\}$, $|S| = preds$, the predicates of S have arity between min and max , $|\text{dom}(D)| = dsize$, and, for each $R \in S$, $|\{\vec{c} \mid R(\vec{c}) \in D\}| = rsizes$. The TGD generator takes as input a set S of predicates and a tuple of values for the tuning parameters ($ssize$, min , max , $tsizes$, $tclass$), and randomly constructs a set Σ of

TGDs such that $\text{sch}(\Sigma) \subseteq S$, $|\text{sch}(\Sigma)| = ssize$, the predicates of $\text{sch}(\Sigma)$ have arity between min and max , $|\Sigma| = tsizes$, and Σ falls in the class $tclass$. We are now ready to proceed with our experimental evaluation. Note that for the experiments we used a server with an Intel Core i5 3.00GHz CPU and 16GB RAM, all the databases in our experiments are stored in a PostgreSQL 11.5 instance, and the termination algorithms have been implemented in Java SE 11.

7 EVALUATION FOR SIMPLE-LINEAR TGDs

We start with the experimental evaluation of IsChaseFinite[SL], which is depicted in Algorithm 1. Towards a refined analysis, we are going to break down its end-to-end runtime, which we denote by $t\text{-total}$, into the following three time parameters:

- $t\text{-parse}$: time to parse the TGDs from an input file,
- $t\text{-graph}$: time to build the dependency graph G of the input set of TGDs, and
- $t\text{-comp}$: time to find the special SCCs in the graph G .

In the rest of the section, we explain how we generate the sets of simple-linear TGDs that are used in our experimental evaluation, and then present our experimental results and discuss the take-home messages. But let us first give a clarification remark.

Remark. In our analysis, we neglect the time taken by the procedure Supports, which, as explained in Section 5, consists of two steps: (1) find the predicates occurring in the input database, and (2) traverse the dependency graph starting from the positions in the special SCCs in the reverse order to reach the positions of the predicates computed in the first step. Step (1) is performed by running a fast query on the catalog of the DBMS storing the input database, and can be safely ignored as it does not impact the rest of the algorithm. Concerning step (2), the time to traverse the dependency graph is negligible compared to the time needed to find the special SCCs, which is already in the order of milliseconds. Therefore, Supports takes insignificant time compared to the rest of the algorithm. Hence, in our experiments, we assume that all the predicates used by the set of simple-linear TGDs occur in the database, and thus, all the positions in the special SCCs are trivially supported. This also simplifies our experiments as they can be conducted using a very simple database that can be induced by the set Σ of simple-linear TGDs, denoted D_Σ , without using our data generator. In fact, D_Σ has an atom $R(c_1, \dots, c_n)$, where $c_1, \dots, c_n$ are distinct constants, for each predicate $R \in \text{sch}(\Sigma)$.

7.1 Generating Simple-Linear TGDs

We now discuss how the sets of simple-linear TGDs used in our experiments are generated. To systematically generate a representative family of sets of TGDs, without favouring any of the two key parameters, namely the size of the underlying schema and the number of TGDs, we consider three *predicate profiles* consisting of sets of TGDs that mention [5,200], [200,400], and [400,600] predicates of arity between 1 and 5, and we further consider three *TGD profiles* consisting of sets of TGDs with [1,333K], [333K,666K], and [666K,1M] TGDs. Note that our choice to fix the arity of the predicates between 1 and 5 is consistent with what we observe in real-life ontologies, where the arity is typically small. The combination of those predicate and TGD profiles gives rise to nine *combined profiles*

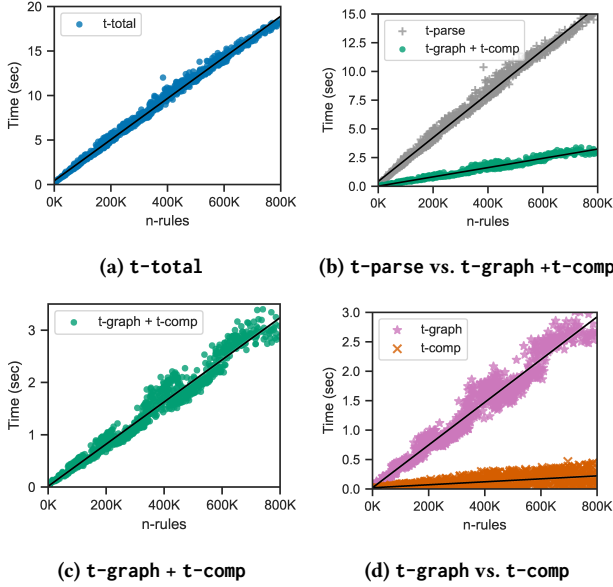


Figure 1: Runtime of IsChaseFinite[SL].

consisting of sets of TGDs with similar syntactic properties. For example, the combined profile obtained from the predicate profile [200,400] and the TGD profile [333K,666K] consists of sets of TGDs Σ such that $200 \leq |\text{sch}(\Sigma)| \leq 400$, each predicate of $\text{sch}(\Sigma)$ has arity between 1 and 5, and $333K \leq |\Sigma| \leq 666K$. For our experiments, we generated 100 sets of TGDs for each of the nine combined profiles as follows. We have first constructed the underlying schema S by generating 1000 predicates, while their arities were randomly selected from [1,5]. Then, for the combined profile induced by the predicate profile $[x, y]$ and the TGD profile $[z, w]$, we have generated 100 sets of simple-linear TGDs by repeatedly executing our TGD generator with input the schema S and the tuple of values for the tuning parameters ($ssize, 1, 5, tsize, SL$), where $ssize$ and $tsize$ were randomly chosen from $[x, y]$ and $[z, w]$, respectively.

7.2 Experimental Evaluation

The algorithm IsChaseFinite[SL] was run for each one of the 900 sets of TGDs of the combined profiles discussed above. The scatter plots in Figure 1 show the runtime of IsChaseFinite[SL](D_Σ, Σ), for each set Σ of TGDs from the combined profiles, i.e., each point in the plots corresponds to one of the 900 sets of TGDs. Figure 1a shows the total runtime ($t\text{-total}$) for sets of TGDs with various sizes ($n\text{-rules}$). Figure 1b breaks down $t\text{-total}$ into the time to parse the TGDs ($t\text{-parse}$) and the time to build their dependency graph and find the special SCCs ($t\text{-graph} + t\text{-comp}$). Figure 1c zooms in $t\text{-graph} + t\text{-comp}$, which are shown separately in Figure 1d.

It is evident from the above scatter plots that the time parameters $t\text{-parse}$ and $t\text{-graph}$ increase linearly as long as we increase $n\text{-rules}$, whereas $t\text{-comp}$ increases very slowly. Let us remark that we have not observed such a linear relationship (in fact, we have not observed any correlation) between the time parameters $t\text{-parse}$ and $t\text{-graph}$, and the number of predicates of the underlying schema. The linear relationship between $t\text{-parse}$ and $n\text{-rules}$

is because parsing each TGD takes constant time since the arity of the predicates falls in the limited range [1,5], and each TGD has one atom in its body and one atom in its head. Note that allowing multi-heads will not change this since, as discussed in Section 6, the number of head-atoms is negligible compared to the number of TGDs. The linear relationship with $t\text{-graph}$ (as shown in Figure 1d) is because the algorithm iterates over the TGDs and spends constant time to process each TGD and update the graph by adding new nodes and edges. Again, since the arity of the predicates falls in [1,5], and each TGD has one atom in its body and one atom in its head, the number of nodes and edges added in the dependency graph due to a certain TGD is in a small fixed range, and thus, the time to update the graph w.r.t. each TGD is constant. The fact that $t\text{-comp}$ increases very slowly is because finding the special SCCs solely depends on the dependency graph, which is in general much smaller than the set of TGDs, while Tarjan’s algorithm is quite efficient that runs in linear time in the size of the underlying graph.

7.3 Take-home Messages

The main takeaway from the experimental results for simple-linear TGDs is that the primary parameter impacting the runtime of IsChaseFinite[SL] is the number of TGDs ($n\text{-rules}$), and we have also observed that the algorithm is very fast even for extremely large sets of TGDs. In fact, most of the end-to-end runtime is spent on parsing ($t\text{-parse}$) and building the dependency graph ($t\text{-graph}$), whereas the time to find special SCCs ($t\text{-comp}$) is insignificant compared to $t\text{-parse}$ and $t\text{-graph}$. To be more precise, $t\text{-parse}$ is much larger than $t\text{-graph}$, and it actually takes most of the total end-to-end runtime of the algorithm. This illustrates the effectiveness of IsChaseFinite[SL] as the actual check for the finiteness of the chase instance for large sets of TGDs is much faster than even reading and parsing the TGDs from the input file.

8 EVALUATION FOR LINEAR TGDs

We now proceed with the evaluation of IsChaseFinite[L], depicted in Algorithm 3. Differently from IsChaseFinite[SL], where the input database did not play any crucial role, we now have a component that heavily relies on the database, that is, the procedure that computes the database shapes, which is part of dynamic simplification. In other words, we have the *database-dependent component* of IsChaseFinite[L], that is, find the database shapes, and the *database-independent component*, that is, simplify the given set of linear TGDs by using the database shapes, build the dependency graph of the simplified set of TGDs, and find the special SCCs in this graph. We claim that these two components, which from now on we call db-dependent and db-independent, respectively, should be evaluated separately as their runtime is impacted by different parameters. Concerning the db-dependent component, it is obvious that it is only affected by the database, whereas the set of TGDs plays no role. On the other hand, although it is clear that the db-independent component is affected by the set of TGDs, it is not straightforward to see that it is not affected by the input database since it operates on a dynamically simplified set of TGDs. Interestingly, we experimentally confirm below that this is indeed the case. Consequently, towards a refined analysis of the algorithm IsChaseFinite[L], we are going to consider the following four time parameters:

- t-shapes: time to find the database shapes,
- t-parse: time to parse the TGDs from an input file,
- t-graph: time to build the dependency graph G of the simplified version of the input set of TGDs (including the time for the simplification using the database shapes), and
- t-comp: time to find the special SCCs in the graph G .

Clearly, t-shapes refers to the runtime of the db-dependent component, whereas t-parse + t-graph + t-comp, which we denote by t-total, refers to the end-to-end runtime of the db-independent component. We proceed to explain how we generate the databases and the sets of linear TGDs that are used in our experimental evaluation, confirm that the db-independent component is not affected by the database, present our experimental results for the two components of IsChaseFinite[L], and discuss the take-home messages.

8.1 Generating Databases and Linear TGDs

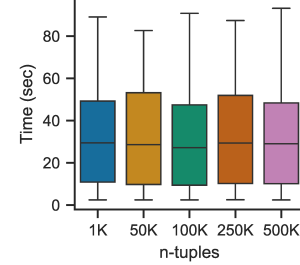
The goal is to generate a family of pairs of the form (D, Σ) , where D is a database and Σ a set of linear TGDs, that will serve as the input to IsChaseFinite[L] in the experimental evaluation. To this end, we first constructed a very large database, which we dub D^* , by using our data generator with input (1000, 1, 5, 500K, 500K). The database D^* mentions 1000 predicates of arity between 1 and 5, and each such predicate has 500K tuples, resulting in a very large database with 500M tuples in total. We further devised views over D^* that allow us to define on-demand virtual databases with 1K, 50K, 100K, 250K, and 500K tuples per predicate, resulting in databases with 1M, 50M, 100M, 250M, and 500M tuples in total, respectively. Those views are actually implemented as a simple SQL query that simply keeps the first 1K, 50K, 100K, 250K, and 500K tuples per predicate, respectively. Note that the tuples in D^* are lexicographically sorted, which means that the different shapes in a relation of D^* are evenly distributed. This in turn ensures that the virtual databases defined via the views have a variety of shapes, which is crucial for our purposes. Having the database D^* and the database views in place the desired family of pairs was generated as described below.

For each one of the nine combined profiles used in the generation of simple-linear TGDs in Section 7, we generated 5 sets of linear TGDs, totalling 45 sets. In particular, for the combined profile induced by the predicate profile $[x, y]$ and the TGD profile $[z, w]$, we have generated 5 sets of linear TGDs by repeatedly executing our TGD generator with input the schema $\{R \mid R(\bar{c}) \in D^*\}$, i.e., the 1000 predicates occurring in D^* , and the tuple of values for the tuning parameters $(ssize, 1, 5, tsize, L)$, where $ssize$ and $tsize$ were randomly chosen from $[x, y]$ and $[z, w]$, respectively. Let Σ^* be the family that collects the 45 generated sets of linear TGDs. Then, for each set $\Sigma \in \Sigma^*$ of linear TGDs, by exploiting the database views discussed above, we obtained five virtual databases of varying size (1K, 50K, 100K, 250K, and 500K tuples per predicate), denoted $D_{\Sigma}^1, D_{\Sigma}^{50}, D_{\Sigma}^{100}, D_{\Sigma}^{250}$, and D_{Σ}^{500} , respectively, leading to five pairs. Summing up, we generated the family $\{(D_{\Sigma}^s, \Sigma) \mid \Sigma \in \Sigma^* \text{ and } s \in \{1, 50, 100, 250, 500\}\}$ consisting of 225 pairs.

8.2 Experimental Evaluation

Before delving into the evaluation of the two components of the algorithm IsChaseFinite[L], let us first confirm that indeed the db-independent component is not affected by the input database.

Separate the Two Components. The figure below depicts the average time over all generated pairs, consisting of a database D_{Σ}^s of a certain size $s \in \{1, 50, 100, 250, 500\}$ and a set Σ of linear TGDs, for building the dependency graph of the dynamically simplified version of Σ using the shapes of D_{Σ}^s and finding the special SCCs:



Interestingly, it confirms that the database size does not impact the time to build the dependency graph and find the special SCCs; thus, it does not impact the end-to-end runtime of the db-independent component, as claimed above. This can be explained by the fact that the number of shapes in a database increases very slowly as we increase the size of the database; this is illustrated in Figure 2. In particular, the bar plots in Figure 2 show the average number of shapes over all databases D_{Σ}^s of a certain size s , where each plot corresponds to a certain predicate profile $[x, y]$, i.e., Σ falls in the predicate profile $[x, y]$. It is clear that the number of shapes increases as we increase the size of the database, which was expected. The interesting outcome, however, is the fact that this increase is very slow, which should be attributed to the fact that many tuples are likely to induce the same shape. Moreover, a new shape gives rise to only a few simplified TGDs, and it does not significantly affect the time for building and processing the dependency graph.

Let us finally observe that, by comparing the three bar plots, it is clear that the number of predicates, reflected in the predicate profile, impacts the number of shapes, which is rather expected as with more predicates there would be more shapes. This means that the number of predicates of the underlying schema is a parameter that affects the number of shapes, which explains why in our analysis above we had to separately consider the three predicate profiles.

Evaluation of the DB-dependent Component. We run the procedure FindShapes for each one of the databases D_{Σ}^s , where $\Sigma \in \Sigma^*$ and $s \in \{1, 50, 100, 250, 500\}$; 225 executions in total. Recall that we have two kinds of implementations for the procedure FindShapes, namely in-memory and in-database (see Section 5). The bar plots in Figure 3 show the average runtime over all databases D_{Σ}^s of a certain size s for finding the shapes in the case of the in-database implementation, where each plot corresponds to a certain predicate profile. Due to space constraints, we do not report the analogous plots for the in-memory implementation. Note, however, that for both implementations we observed a similar trend, with the in-database implementation outperforming the in-memory one.

It is evident from the bar plots in Figure 3 that the time to find the shapes increases while the database size increases, which is not surprising since, as discussed above, the number of shapes increases while the database size increases. Observe, however, that the time to find the shapes grows much faster than the actual number shapes, which should be attributed to the fact that for finding the shapes we

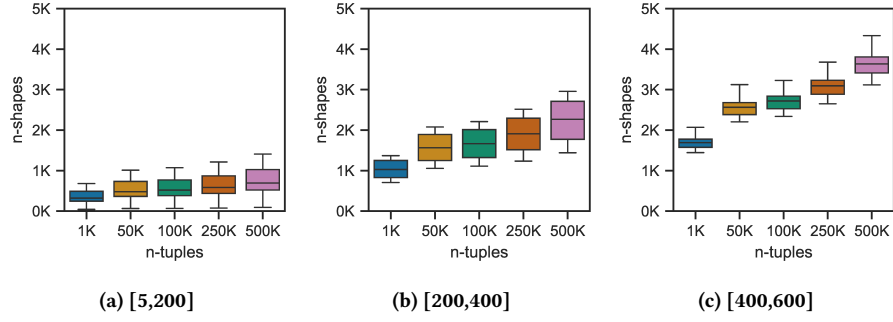


Figure 2: Number of Shapes.

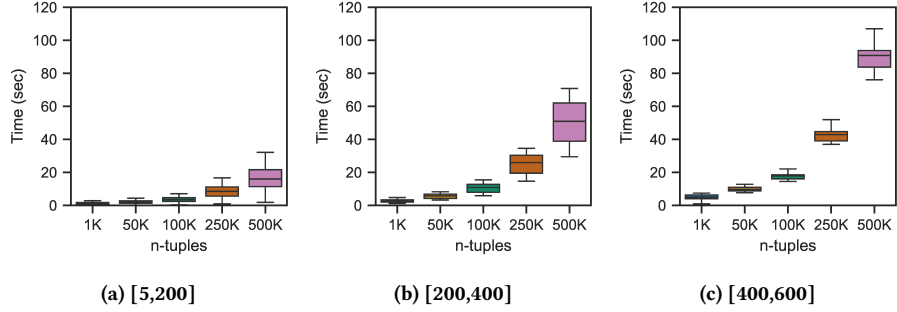


Figure 3: Runtime of FindShapes (in-database implementation).

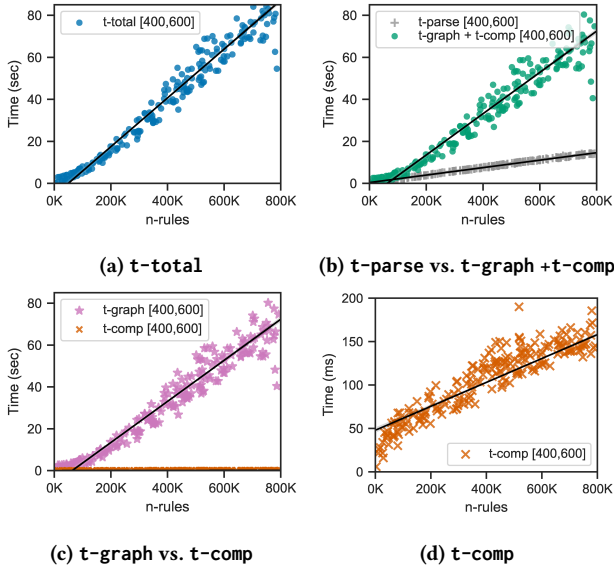


Figure 4: Runtime of the db-independent component.

actually need to scan the whole database. Let us finally observe that, by comparing the three bar plots, it is apparent that the number of predicates, reflected in the underlying predicate profile, also impacts the time to find the shapes, which explains why we had to analyze each predicate profile separately.

Evaluation of the DB-independent Component. The scatter plots in Figure 4 show the runtime of the db-independent component of the algorithm `IsChaseFinite[L]` when executed with input (D_Σ^s, Σ) , for each set $\Sigma \in \Sigma^\star$ falling in the predicate profile [400,600] and $s \in \{1, 50, 100, 250, 500\}$. In particular, each point in the plots corresponds to a pair (D_Σ^s, Σ) . Let us stress that, unlike the analogous Figure 1 for simple-linear TGDs, we focus on a particular predicate profile since otherwise we do not obtain any trend of the runtime w.r.t. the number of TGDs. In other words, the apparent linear trend observed in those plots only holds for sets of TGDs from the same predicate profile. This is because the number of predicates of the underlying schema impacts the number of shapes, which in turn affects the process of dynamic simplification and the size of the dependency graph, and thus, the time parameters $t\text{-graph}$ and $t\text{-comp}$ are impacted. This explains why we had to analyze each predicate profile separately. Due to space constraints, we omit the analogous plots for the smaller predicate profiles [5,200] and [200,400], which depict similar trends. Figure 4a shows the total runtime ($t\text{-total}$) for sets of TGDs with various sizes ($n\text{-rules}$). Figure 4b breaks down $t\text{-total}$ into the time to parse the TGDs ($t\text{-parse}$) and the time to build their dependency graph and find the special SCCs ($t\text{-graph} + t\text{-comp}$), whereas Figure 4c shows separately $t\text{-graph}$ and $t\text{-comp}$. Figure 4d zooms in $t\text{-comp}$.

It is evident from the above scatter plots that the time parameters $t\text{-parse}$ and $t\text{-graph}$ increase linearly as long as we increase $n\text{-rules}$, whereas $t\text{-comp}$ increases very slowly. This is essentially what we have observed for simple-linear TGDs in Figure 1, with the key difference that the time needed to parse the TGDs ($t\text{-parse}$) is now much less compared to the time for building the dependency

graph and finding the special SCCs ($t\text{-graph} + t\text{-comp}$). It should not be forgotten, however, that for linear TGDs we need to focus on a single predicate profile in order to get these linear trends; otherwise, if we consider all the predicate profiles at once, there is no trend that can be observed. Note also that the absolute running time increases compared to the case of simple-linear TGDs.

8.3 Take-home Messages

The main takeaway is that the algorithm `IsChaseFinite[L]` consists of two components that are of different nature in the sense that their runtime is impacted by different parameters. On the one hand, we have the db-dependent component that is only affected by the size of the database. On the other hand, we have the db-independent component whose runtime is primarily affected by the number of TGDs ($n\text{-rules}$). Having said that, we have also observed that the number of predicates also affects the runtime of the db-independent component since it impacts the number of shapes, which in turn affects the process of dynamic simplification and the size of the dependency graph. We conclude by observing that the total runtime of `IsChaseFinite[L]` is quite reasonable, which should be seen as a strong evidence that fast checking for the finiteness of the chase instance in the case of linear TGDs is not an unrealistic goal. Note that most of the total end-to-end runtime of the algorithm is spent on finding database shapes, which indicates that our future efforts should be concentrated on improving the db-dependent component.

9 VALIDATION OF RESULTS

With the aim of validating the main outcome of the stress test analysis performed in Sections 7 and 8, we have also run experiments using databases and sets of TGDs that are available in the literature.

Adopted Scenarios. We considered three families of databases and sets of TGDs from the literature: (i) the family Deep from [4] that collects sets of simple-linear TGDs that are at the same time weakly-acyclic, which has been developed to test data exchange scenarios, (ii) LUBM, which is a popular benchmark consisting of an ontology modelled using the central Description Logic (DL) EL, called Univ-Bench, and a data generator, called UBA, for generating synthetic data over the vocabulary of Univ-Bench [12], and (iii) iBench, which is a framework for generating TGDs with tuning parameters that can control a wide range of properties [2].

Discussion. Concerning `IsChaseFinite[SL]` for simple-linear TGDs, we observed that it runs in a few milliseconds for all the scenarios discussed above. This is a confirmation that for simple-linear TGDs, checking for the finiteness of the chase can be done very efficiently. We have also seen that the validation analysis confirms the main outcome of the analysis for the algorithm `IsChaseFinite[L]` performed in Section 8. In particular, we observed that indeed the costly task is finding the database shapes, whereas the time taken by the db-independent component is negligible. Moreover, we have seen that checking for the finiteness of the chase instance can be done rather efficiently in practice. In particular, for sets consisting of thousands of TGDs such as the Deep scenarios, and millions of facts such as some members of LUBM, it takes less than a second. Another interesting takeaway from the validation analysis is that there is no clear way to go regarding the implementation

of FindShape among the in-memory and the in-database options. In particular, the in-memory implementation is preferred when there are a few tuples per relation in the database, whereas the in-database implementation performs better when the underlying schema has a few predicates of small arity. For schemas with many predicates, each of which has many tuples in the input database, both implementations require significant time, and an offline computation of the database shapes might be preferred.

Feasibility in the Real World. We have also considered real-world scenarios with the aim of analysing the feasibility of our algorithms in a more realistic context. We considered 483 pairs (D, Σ) , where D is a database and Σ a set of simple-linear TGDs, obtained from the ontology repository of the Data and Knowledge Group of the University of Oxford by keeping from the available ontologies only the ontological axioms that can be expressed as simple-linear TGDs. Let us stress that these are real-world ontologies modelled using Description Logics, and include, among others, a large subset of the Gardiner ontology corpus, several Phenoscape ontologies, and a number of ontologies from two versions of the Open Biomedical Ontology (OBO) corpus [11]. Both `IsChaseFinite[SL]` and `IsChaseFinite[L]` performed swiftly with the maximum time taken to be 1.77 and 0.44 seconds, respectively. On average, `IsChaseFinite[SL]` and `IsChaseFinite[L]` took approximately 0.85 and 0.22 seconds, respectively, across different scenarios. A rather unexpected finding is that `IsChaseFinite[L]`, which employs dynamic simplification, outperformed `IsChaseFinite[SL]`. This should be attributed to two reasons: (i) the underlying schema consists of only unary and binary predicates, which in turn allows for a fast computation of shapes, and (ii) dynamic simplification uses a small number of simple-linear TGDs for building the dependency graph, which is a consequence of the fact that the databases mention a small number of predicates from the underlying schema.

10 CONCLUSIONS AND FUTURE WORK

Our work provides the first systematic attempt to experimentally evaluate termination algorithms for the semi-oblivious chase. Our analysis revealed that for simple-linear TGDs, we can efficiently check whether the chase terminates even for very large databases and sets of TGDs. Concerning linear TGDs, the overall runtime of the algorithm is quite reasonable, but there is still room for improvement. Interestingly, our analysis showed that the algorithm for linear TGDs consists of two separate components, the db-dependent and the db-independent ones. This modular nature of the algorithm allows us to study and improve the two components separately. In particular, we have observed that the heavy component is the db-dependent one, and thus, we can focus our future efforts to improve the performance of that component. Although our analysis relied on an in-database and an in-memory implementation of the procedure for finding the shapes, we could adopt other techniques without affecting the db-independent component.

ACKNOWLEDGMENTS

We thank the referees for their feedback. Andreas Pieris was supported by the EPSRC grant EP/S003800/1 "EQUID". Mostafa Milani was supported by the NSERC Discovery Grant 2021-04120.

REFERENCES

- [1] Alfred V. Aho, Yehoshua Sagiv, and Jeffrey D. Ullman. 1979. Efficient Optimization of a Class of Relational Expressions. *ACM Trans. Database Syst.* 4, 4 (1979), 435–454.
- [2] Patricia C. Arocena, Boris Glavic, Radu Ciucanu, and Renée J. Miller. 2015. The iBench Integration Metadata Generator. *PVLDB* 9, 3 (2015), 108–119.
- [3] Catriel Beeri and Moshe Y. Vardi. 1984. A Proof Procedure for Data Dependencies. *J. ACM* 31, 4 (1984), 718–741.
- [4] Michael Benedikt, George Konstantinidis, Giansalvatore Mecca, Boris Motik, Paolo Papotti, Donatello Santoro, and Efthymia Tsamoura. 2017. Benchmarking the Chase. In *PODS*. 37–52.
- [5] Marco Calautti, Georg Gottlob, and Andreas Pieris. 2022. Non-Uniformly Terminating Chase: Size and Complexity. In *PODS*. 369–378.
- [6] Marco Calautti and Andreas Pieris. 2021. Semi-Oblivious Chase Termination: The Sticky Case. *Theory Comput. Syst.* 65, 1 (2021), 84–121.
- [7] Andrea Cali, Georg Gottlob, and Thomas Lukasiewicz. 2012. A general Datalog-based framework for tractable query answering over ontologies. *J. Web Sem.* 14 (2012), 57–83.
- [8] Diego Calvanese, Giuseppe De Giacomo, Domenico Lembo, Maurizio Lenzerini, and Riccardo Rosati. 2007. Tractable Reasoning and Efficient Query Answering in Description Logics: The DL-Lite Family. *J. Autom. Reasoning* 39, 3 (2007), 385–429.
- [9] Alin Deutsch, Alan Nash, and Jeff B. Remmel. 2008. The Chase Revisited. In *PODS*. 149–158.
- [10] Ronald Fagin, Phokion G. Kolaitis, Renée J. Miller, and Lucian Popa. 2005. Data exchange: semantics and query answering. *Theor. Comput. Sci.* 336, 1 (2005), 89–124.
- [11] Bernardo Cuenca Grau, Ian Horrocks, Markus Krötzsch, Clemens Kupke, Despoina Magka, Boris Motik, and Zhe Wang. 2013. Acyclicity Notions for Existential Rules and Their Application to Query Answering in Ontologies. *J. Artif. Intell. Res.* 47 (2013), 741–808.
- [12] Yuanbo Guo, Zhengxiang Pan, and Jeff Heflin. 2005. LUBM: A benchmark for OWL knowledge base systems. *J. Web Semant.* 3, 2-3 (2005), 158–182.
- [13] Markus Krötzsch, Maximilian Marx, and Sebastian Rudolph. 2019. The Power of the Terminating Chase (Invited Talk). In *ICDT*. 3:1–3:17.
- [14] Michel Leclère, Marie-Laure, Michaël Thomazo, and Federico Ulliana. 2019. A Single Approach to Decide Chase Termination on Linear Existential Rules. In *ICDT*. 18:1–18:19.
- [15] David Maier, Alberto O. Mendelzon, and Yehoshua Sagiv. 1979. Testing Implications of Data Dependencies. *ACM Trans. Database Syst.* 4, 4 (1979), 455–469.
- [16] Bruno Marnette. 2009. Generalized schema-mappings: from termination to tractability. In *PODS*. 13–22.
- [17] Yavor Nenov, Robert Piro, Boris Motik, Ian Horrocks, Zhe Wu, and Jay Banerjee. 2015. RDFox: A Highly-Scalable RDF Store. In *ISWC*. 3–20.
- [18] Robert Endre Tarjan. 1972. Depth-First Search and Linear Graph Algorithms. *SIAM J. Comput.* 1, 2 (1972), 146–160.
- [19] Jacopo Urbani, Markus Krötzsch, Criel J. H. Jacobs, Irina Dragoste, and David Carral. 2018. Efficient Model Construction for Horn Logic with VLog - System Description. In *IJCAR*. 680–688.