

Pay-as-you-go Ontology Query Answering Using a Datalog Reasoner^{*}

Yujiao Zhou, Yavor Nenov, Bernardo Cuenca Grau, and Ian Horrocks

Department of Computer Science, University of Oxford, UK

Abstract. We describe a hybrid approach to conjunctive query answering over OWL 2 ontologies that combines a datalog reasoner with a fully-fledged OWL 2 reasoner in order to provide scalable “pay as you go” performance. Our approach delegates the bulk of the computation to the highly scalable datalog engine and resorts to expensive OWL 2 reasoning only as necessary to fully answer the query. We have implemented a prototype system that uses RDFS as a datalog reasoner, and HermiT as an OWL 2 reasoner. Our evaluation over both benchmark and realistic ontologies and datasets suggests the feasibility of our approach.

1 Introduction

The use of RDF [14], OWL 2 [15], and SPARQL 1.1 [24] to represent and query semi-structured data together with domain knowledge is increasingly widespread. Query answering in this setting is, however, of high worst-case complexity [6, 5], and although heavily optimised, existing systems for query answering w.r.t. RDF data and an unrestricted OWL 2 ontology can process only small to medium size datasets [12, 25, 10]. This has led to the development of query answering procedures that are more scalable, but that can (fully) process only fragments of OWL 2, and several prominent fragments have now been standardised as OWL 2 profiles [13]. Such systems have been shown to be (potentially) highly scalable [23, 1, 22, 26], but if the ontology falls outside the relevant profile, then the answers computed by such a system may be incomplete: if it returns an answer, then all tuples in the answer are (usually) valid, but some valid tuples may be missing from the answer. When used with out-of-profile ontologies, a query answer computed by such a system can thus be understood as providing a *lower-bound* on the correct answer; however, they cannot in general provide any upper bound or even any indication as to how complete the computed answer is [2].

In this paper, we describe a hybrid approach to query answering that exploits a datalog reasoner to compute both a lower bound answer and an upper bound answer. If lower and upper bound answers coincide, they obviously provide a sound and complete answer. Otherwise, *relevant fragments* of the ontology and data can be extracted that are guaranteed to be sufficient to test the validity of tuples in the “gap” between the two answers. These fragments can also be

^{*} This paper recapitulates some of our results in [27–29], and it is accompanied with a technical report available at <http://www.cs.ox.ac.uk/isg/people/yujiao.zhou>.

computed by relying solely on the datalog reasoner, and are typically much smaller than the input ontology and data. The remaining gap tuples need to be checked w.r.t. to the identified fragments using an OWL 2 reasoner such as HermiT [16] or Pellet [18]; furthermore, since the number of gap tuples can be significant in some cases, we exploit summarisation techniques inspired by the SHER system [3, 4] to quickly identify spurious gap answers, thus further reducing the requirement for fully-fledged OWL 2 reasoning.

Our approach is *pay-as-you-go* in the sense that the bulk of the computation is delegated to a scalable datalog engine. Furthermore, although our main goal is to answer queries over OWL 2 ontologies efficiently, our technical results are very general and our approach is not restricted to DLs. More precisely, given a first-order KR language L that can be captured by rules allowing for existential quantification and disjunction in the head, and over which we want to answer conjunctive queries, our only assumption is the availability of a fully-fledged reasoner for L and a datalog reasoner, which are both used as a “black box”.

We have implemented our techniques in a prototypical system using the RDFox as a datalog reasoner [23] and the HermiT as a fully-fledged OWL 2 reasoner.¹ Our preliminary evaluation over both benchmark and realistic data suggests that the system can provide scalable pay-as-you-go query answering for a wide range of OWL 2 ontologies, RDF data and queries. In almost all cases, the system is able to completely answer queries without resorting to fully-fledged OWL 2 reasoning, and even when this is not the case, relevant fragment extraction and summarisation are effective in reducing the size of the problem to manageable proportions.

2 Preliminaries

We adopt standard first order logic notions, such as variables, constants, atoms, formulas, clauses, substitutions, satisfiability, and entailment. We also assume basic familiarity with OWL 2 [15] and its profiles [13]. A generalised rule (or just a *rule*) is a function-free sentence of the form

$$\forall \mathbf{x} \left(\bigwedge_{j=0}^n B_j(\mathbf{x}) \rightarrow \bigvee_{i=0}^m \exists \mathbf{y}_i \varphi_i(\mathbf{x}, \mathbf{y}_i) \right)$$

where $B_j(\mathbf{x})$ are *body* atoms and φ_i are conjunctions of *head* atoms. The universal quantifiers are left implicit from now on. A rule is *Horn* if $m \leq 1$, and it is *datalog* if it is Horn and does not contain existential quantifiers. A *fact* is a ground atom and a *dataset* is a finite set of facts. A *knowledge base* $\mathcal{K}$ consists of a finite set of rules and a dataset. We treat equality ($\approx$) as an ordinary predicate, but assume that every knowledge base in which equality occurs contains the axioms of equality for its signature. Each OWL 2 ontology can be normalised as one

¹ Although our techniques are proved correct for general conjunctive queries, in practice we are limited by the current query capabilities of OWL 2 reasoners.

$\text{Foreman}(x) \rightarrow \text{Manag}(x)$	(T_1)
$\text{Superv}(x) \rightarrow \text{Manag}(x)$	(T_2)
$\text{Superv}(x) \wedge \text{boss}(x, y) \rightarrow \text{Workman}(y)$	(T_3)
$\text{TeamLead}(x) \wedge \text{boss}(x, y) \wedge \text{Manag}(y) \rightarrow$	(T_4)
$\text{Manag}(x) \rightarrow \text{Superv}(x) \vee \exists y.(\text{boss}(x, y) \wedge \text{Manag}(y))$	(T_5)
$\text{Manag}(x) \rightarrow \exists y.(\text{boss}(x, y))$	(T_6)

$\text{Manag}(\text{Sue})$	(D_1)	$\text{Superv}(\text{Rob})$	(D_3)	$\text{Manag}(\text{Jo})$	(D_5)
$\text{Superv}(\text{Dan})$	(D_2)	$\text{boss}(\text{Dan}, \text{Ben})$	(D_4)	$\text{TeamLead}(\text{Jo})$	(D_6)
				$\text{boss}(\text{Jane}, \text{Rob})$	(D_7)

Fig. 1. Example knowledge base $\mathcal{K}^{ex}$.

such knowledge base using the correspondence of OWL and first order logic and a variant of the structural transformation (e.g., see [16] for details).

We focus on CQ answering as the key reasoning problem. A *query* is a formula $q(\mathbf{x}) = \exists \mathbf{y} \varphi(\mathbf{x}, \mathbf{y})$ with $\varphi(\mathbf{x}, \mathbf{y})$ a conjunction of atoms. We usually omit the free variables $\mathbf{x}$ of queries and write just q . The query is *atomic* if $\varphi(\mathbf{x}, \mathbf{y})$ is a single atom. A tuple of individuals $\mathbf{a}$ is a (*certain*) *answer* to q w.r.t. a set of sentences $\mathcal{F}$ iff $\mathcal{F} \models q(\mathbf{a})$. The set of all answers to $q(\mathbf{x})$ w.r.t. $\mathcal{F}$ is denoted by $\text{cert}(q, \mathcal{F})$.

There are two main techniques for answering queries over a datalog knowledge base $\mathcal{K}$. *Forward chaining* computes the set $\text{Mat}(\mathcal{K})$ of ground atoms entailed by $\mathcal{K}$, called the *materialisation* of $\mathcal{K}$. A query q over $\mathcal{K}$ can be answered directly over the materialisation. *Backward chaining* treats a query as a conjunction of atoms (a *goal*). An *SLD resolvent* of a goal $A \wedge \psi$ with a datalog rule $\varphi \rightarrow C_1 \wedge \dots \wedge C_n$ is a goal $\psi\theta \wedge \varphi\theta$, with θ the MGU of A and C_j , for some $1 \leq j \leq n$. An *SLD proof* of a goal G_0 in $\mathcal{K}$ is a sequence of goals $(G_0, \dots, G_n)$ with G_n the empty goal ($\square$), and each G_{i+1} a resolvent of G_i and a rule in $\mathcal{K}$.

3 Overview

The main idea behind our approach to query answering is to delegate the bulk of the computational workload to a highly scalable datalog reasoner, thus minimising the use of a fully-fledged OWL 2 reasoner. Given a knowledge base $\mathcal{K}$ and a query q , we proceed according to the following algorithm:

1. Use the datalog reasoner to compute both lower bound (sound but possibly incomplete) and upper bound (complete but possibly unsound) answers to the (Boolean) unsatisfiability query and the input q . (See Sections 4, 5).
2. If both bounds report unsatisfiability, then we return unsatisfiable. If none of them reports unsatisfiability and they yield the same answers to q , we output the resulting answers. In any other case, proceed to the next step.

3. Use the datalog reasoner to compute fragments $\mathcal{K}_\perp$ and $\mathcal{K}_{[q,G]}$ of $\mathcal{K}$, where G is the set of answers to q in the gap between the bounds. (See Section 6).
4. If the upper bound reports unsatisfiability and $\mathcal{K}_\perp$ is unsatisfiable, then return unsatisfiable.
5. Use the OWL reasoner to check whether $\mathcal{K}_{[q,G]} \cup \mathcal{K}_\perp \models q(\mathbf{a})$, for each $\mathbf{a} \in G$. To minimise the computational workload of the OWL reasoner, this step is carried out as follows (see Section 7):
 - (a) Summarise $\mathcal{K}_{[q,G]} \cup \mathcal{K}_\perp$ by merging all constants that instantiate the same unary predicates [4]. Use the OWL reasoner to discard those $\mathbf{a} \in G$ such that $q(\mathbf{a})$ is not entailed by the summarised KB.
 - (b) Compute a dependency relation between the remaining elements of G such that if $\mathbf{b}$ depends on $\mathbf{a}$ and $\mathbf{a}$ is a spurious answer, then so is $\mathbf{b}$. Arrange the calls to the reasoners according to these dependencies.
6. Return the lower bound answers to q plus those tuples in G determined to be answers in Step 5.

We will describe each of these steps and illustrate them using as running example the knowledge base $\mathcal{K}^{ex}$ in Figure 1 and the following query q^{ex} :

$$q^{ex}(x) = \exists y(\text{boss}(x, y) \wedge \text{Workman}(y))$$

4 Computing Upper Bounds

To compute upper bound query answers, we first compute a datalog knowledge base $\mathcal{U}(\mathcal{K})$ that entails the nullary predicate $\perp$, if $\mathcal{K}$ is unsatisfiable, and that entails $\mathcal{K}$, otherwise. Hence, for satisfiable knowledge bases $\mathcal{K}$ we get that $\text{cert}(q, \mathcal{U}(\mathcal{K}))$ subsumes $\text{cert}(q, \mathcal{K})$. The knowledge base $\mathcal{U}(\mathcal{K})$ is the result of consecutively applying the transformations Σ , Ξ and Ψ defined next.

Definition 1. Let $\mathcal{K}$ be a KB. We define $\mathcal{U}(\mathcal{K}) := \Psi \circ \Xi \circ \Sigma(\mathcal{K})$, where

- Σ is a mapping that transforms each rule into clausal normal form;
- Ξ maps each clause C to a set of clauses as follows: (i) if C contains only negative literals, then $\Xi(C) = C \vee \perp$; (ii) if C is of the form $\neg B_0 \vee \dots \vee \neg B_k \vee C_0 \vee \dots \vee C_{r+1}$ then $\Xi(C)$ consists of the clauses $\neg B_1 \vee \dots \vee \neg B_k \vee C_i$, for $0 \leq i \leq r+1$; (iii) in any other case, $\Xi(C) = C$.
- Ψ maps every Horn clause C to a datalog rule $\Psi(C)$ obtained from C by first replacing each functional term with a globally fresh constant, and then transforming the resulting clause into its equivalent datalog rule.

These transformations extend to sets in the natural way.

In our example $\mathcal{K}^{ex}$, the transformation $\mathcal{U}$ is the identity for all rules except T_4 – T_6 . Rule T_4 is transformed by $\mathcal{U}$ (Ξ in particular) into the datalog rule U_4 .

$$\text{TeamLead}(x) \wedge \text{boss}(x, y) \wedge \text{Manag}(y) \rightarrow \perp \quad (U_4)$$

Rule T_5 is first transformed by Σ into clauses $\neg\text{Manag}(x) \vee \text{Superv}(x) \vee \text{boss}(x, f_1(x))$ and $\neg\text{Manag}(x) \vee \text{Superv}(x) \vee \text{Manag}(f_1(x))$. These will then be transformed by Ξ to the clauses $\neg\text{Manag}(x) \vee \text{Superv}(x)$, $\neg\text{Manag}(x) \vee \text{boss}(x, f_1(x))$ and $\neg\text{Manag}(x) \vee \text{Manag}(f_1(x))$. Finally, Ψ will produce the following datalog rules.

$$\begin{aligned} \text{Manag}(x) &\rightarrow \text{Superv}(x) & (U_5^1) \\ \text{Manag}(x) &\rightarrow \text{boss}(x, c_1) & (U_5^2) \\ \text{Manag}(x) &\rightarrow \text{Manag}(c_1) & (U_5^3) \end{aligned}$$

Rule T_6 will be transformed by Σ into the clause $\neg\text{Manag}(x) \vee \text{boss}(x, f_2(x))$, which in turn will be transformed by Ψ into the datalog rule U_6 .

$$\text{Manag}(x) \rightarrow \text{boss}(x, c_2) \quad (U_6)$$

Hence, $\mathcal{U}(\mathcal{K})$ comprises the facts D_1 – D_7 and the datalog rules T_1 – T_3 , U_4 , U_5^1 – U_5^3 , and U_6 . One can easily verify that $\text{cert}(q^{ex}, \mathcal{U}(\mathcal{K}^{ex})) = \{Sue, Dan, Rob, Jo\}$.

The following lemma captures the properties of these transformations.

Proposition 1. *Let $\mathcal{K}$ be a knowledge base and q be a query. Then:*

1. $\mathcal{K}$ *unsatisfiable* $\Leftrightarrow \Sigma(\mathcal{K})$ *unsatisfiable* $\Rightarrow \Xi(\Sigma(\mathcal{K})) \models \perp \Rightarrow \mathcal{U}(\mathcal{K}) \models \perp$;
2. $\mathcal{K}$ *satisf.* $\Rightarrow \text{cert}(q, \mathcal{K}) = \text{cert}(q, \Sigma(\mathcal{K})) \subseteq \text{cert}(q, \Xi(\Sigma(\mathcal{K}))) \subseteq \text{cert}(q, \mathcal{U}(\mathcal{K}))$.

5 Computing Lower Bounds

A direct way to compute lower bound query answers given $\mathcal{K}$ and q is to select the datalog fragment $\mathcal{L}(\mathcal{K})$ of $\mathcal{K}$, check its satisfiability, and compute $\text{cert}(q, \mathcal{L}(\mathcal{K}))$ using a datalog engine. By monotonicity of first-order logic, $\mathcal{K}$ entails $\mathcal{L}(\mathcal{K})$, and hence $\text{cert}(q, \mathcal{K})$ is guaranteed to subsume $\text{cert}(q, \mathcal{L}(\mathcal{K}))$. In our running example, the lower bound knowledge base $\mathcal{L}(\mathcal{K}^{ex})$ comprises the facts D_1 – D_7 and the datalog rules T_1 – T_4 , and it can be easily verified that $\text{cert}(q^{ex}, \mathcal{L}(\mathcal{K}^{ex})) = \{Dan\}$.

To improve this bound, we adopt the *combined approach* introduced to handle query answering in $\mathcal{ELHO}_\perp^r$ [19, 11]. Given an $\mathcal{ELHO}_\perp^r$ knowledge base $\mathcal{K}'$ and a query q , the combined approach first exploits the upper bound datalog program $\mathcal{U}(\mathcal{K}')$ to check satisfiability of $\mathcal{K}'$ and to compute $\text{cert}(q, \mathcal{U}(\mathcal{K}'))$. A subsequent filtering step Φ , which is efficiently implementable, guarantees to eliminate all spurious tuples; the resulting answer $\Phi(\text{cert}(q, \mathcal{U}(\mathcal{K}')))$ is thus sound and complete w.r.t. q and $\mathcal{K}'$.

The combined approach is clearly compatible with ours. Given an OWL 2 knowledge base $\mathcal{K}$ and query q , we proceed as follows. First, we select the datalog fragment $\mathcal{K}_1 = \mathcal{L}(\mathcal{K})$, and compute the materialisation $\text{Mat}(\mathcal{K}_1)$ using the datalog engine. Second, we select the subset $\mathcal{K}_2$ of $\mathcal{K}$ corresponding to $\mathcal{ELHO}_\perp^r$ axioms and Skolemise existential quantifiers to constants to obtain $\mathcal{U}(\mathcal{K}_2)$. Then, we further compute the answers $\text{cert}(q, \mathcal{U}(\mathcal{K}_2) \cup \text{Mat}(\mathcal{K}_1))$. Finally, we apply the filtering step Φ to obtain the final set of lower bound answers. The $\mathcal{ELHO}_\perp^r$ fragment for our running example $\mathcal{K}^{ex}$ consists of rules T_1 – T_4 and T_6 , and the resulting new lower bound answer of q^{ex} is the set $\{Dan, Rob\}$.

Table 1. SLD proofs of $\perp$ and $q^{ex}(Jo)$ in $\mathcal{U}(\mathcal{K}^{ex})$

$\perp$		$b(J, y) \wedge W(y)$	
$T(x) \wedge b(x, y) \wedge M(y)$	by U_4	$M(J) \wedge W(c_2)$	by U_6
$b(J, y) \wedge M(y)$	by D_6	$W(c_2)$	by D_5
$M(J) \wedge M(c_1)$	by U_5^2	$S(x) \wedge b(x, c_2)$	by T_3
$M(c_1)$	by D_5	$M(x) \wedge b(x, c_2)$	by U_5^1
$M(x)$	by U_3^3	$b(J, c_2)$	by D_5
$\square$	by D_5	$M(J)$	by U_6
		$\square$	by D_5

6 Computing Relevant Fragments

Fragment Definition and Formal Properties The relevant fragments $\mathcal{K}_\perp$ and $\mathcal{K}_{[q, G]}$ are defined in terms of SLD proofs in $\mathcal{U}(\mathcal{K})$. In particular, $\mathcal{K}_\perp$ is defined in terms of proofs for the nullary predicate $\perp$, and $\mathcal{K}_{[q, G]}$ is defined in terms of proofs for each answer in G .

Definition 2. Let $\mathcal{K}$ be a knowledge base, $q(\mathbf{x})$ be a query, and S be a set of tuples. Then $\mathcal{K}_\perp$ (resp. $\mathcal{K}_{[q, S]}$) is the set of all $\alpha \in \mathcal{K}$ for which there exists $\beta \in \mathcal{U}(\alpha)$ involved in an SLD proof of $\perp$ (resp. $Q(\mathbf{a})$, for some $\mathbf{a} \in S$) in $\mathcal{U}(\mathcal{K})$.

The properties of these fragments needed to ensure the correctness of our algorithm in Section 3 are summarised in the following theorem.

Theorem 1. Let $\mathcal{K}$ be a knowledge base, $q(\mathbf{x})$ a conjunctive query, and S a set of tuples. Then, (i) $\mathcal{K}$ is satisfiable iff $\mathcal{K}_\perp$ is satisfiable; and (ii) if $\mathcal{K}$ is satisfiable, then $\mathcal{K} \models q(\mathbf{a})$ iff $\mathcal{K}_{[q, S]} \cup \mathcal{K}_\perp \models q(\mathbf{a})$ for every $\mathbf{a} \in S$.

As expected, $\mathcal{K}_\perp$ can be used to determine satisfiability of $\mathcal{K}$. In case $\mathcal{K}$ is found satisfiable, the union of $\mathcal{K}_{[q, G]}$ and $\mathcal{K}_\perp$ can then be used to check the validity of each candidate answer in G (in this case, $\mathcal{K}_\perp$ is still needed to account for the possible interactions between non-Horn rules and rules with empty heads).

Table 1 specifies proofs of $\perp$ and $q^{ex}(Jo)$ in $\mathcal{U}(\mathcal{K}^{ex})$, where predicates and constants are abbreviated to their first letters. By Definition 2, $\mathcal{K}_\perp \cup \mathcal{K}_{[q^{ex}, \{Jo\}]}$ subsumes $\{T_3, \dots, T_6, D_5, D_6\}$, and, hence, it entails $q^{ex}(Jo)$, as expected. Note that $\mathcal{K}_{[q, \{Jo\}]}$ alone is not sufficient to show $q^{ex}(Jo)$ since every fragment of $\mathcal{K}^{ex}$ that entails $q^{ex}(Jo)$ must include rule T_4 . According to Definition 1, $\mathcal{K}_{[q, \{Jo\}]}$ will include T_4 if and only if U_4 is used in an SLD proof of $q^{ex}(Jo)$ in $\mathcal{U}(\mathcal{K}^{ex})$; however, no such proof will involve U_4 since the goal $q^{ex}(Jo)$ does not involve $\perp$, and there is no way of eliminating $\perp$ from a goal using the rules in $\mathcal{U}(\mathcal{K}^{ex})$ as they do not contain $\perp$ in their bodies.

The proof of Theorem 1 is involved, and details are deferred to the appendix. Nonetheless, we next sketch the arguments behind the proof. A first observation is that, w.l.o.g. we can restrict ourselves to the case where $q(\mathbf{x})$ is atomic.

Lemma 1. Let $\mathcal{K}$ be a knowledge base, $q(\mathbf{x}) = \exists \mathbf{y} \varphi(\mathbf{x}, \mathbf{y})$ be a CQ, S be a set of tuples, Q be a fresh predicate, and let $\mathcal{K}' = \mathcal{K}_{[q, S]} \cup \mathcal{K}_\perp$. Then, $\mathcal{K}' \models q(\mathbf{a})$ iff $\mathcal{K}' \cup \{\varphi(\mathbf{x}, \mathbf{y}) \rightarrow Q(\mathbf{x})\} \models Q(\mathbf{a})$.

The crux of the proof relies on the following properties of Ξ (the step in the definition of $\mathcal{U}$ which splits each non-Horn clause C into Horn clauses).

Lemma 2. *Let $\mathcal{N}$ be a set of first-order clauses. Then:*

- *if $C \in \mathcal{N}$ participates in a refutation in $\mathcal{N}$, then every $C' \in \Xi(C)$ is part of an SLD proof of $\perp$ in $\Xi(\mathcal{N})$;*
- *if $C \in \mathcal{N}$ participates in a resolution proof in $\mathcal{N}$ of an atomic query $Q(\mathbf{a})$, then each $C' \in \Xi(C)$ participates in an SLD proof of $\perp$ or $Q(\mathbf{a})$ in $\Xi(\mathcal{N})$.*

Thus, by Lemma 2, each resolution proof in a set of clauses $\mathcal{N}$ can be mapped to SLD proofs in $\Xi(\mathcal{N})$ that “preserves” the participating clauses. The following lemma allows us to restate Lemma 2 for $\Psi \circ \Xi$ instead of Ξ .

Lemma 3. *Let $\mathcal{H}$ be a set of first-order Horn clauses, $Q(\mathbf{x})$ be an atomic query, and $\mathbf{a}$ be a tuple of constants. If a clause C participates in an SLD proof of $Q(\mathbf{a})$ in $\mathcal{H}$, then $\Psi(C)$ participates in an SLD proof of $Q(\mathbf{a})$ in $\Psi(\mathcal{H})$.*

With these Lemmas, we can exploit refutational completeness of resolution and the entailment preservation properties of Skolemisation to show Theorem 1.

Fragment Computation The computation of the relevant fragments requires a scalable algorithm for “tracking” all rules and facts involved in SLD proofs for datalog programs. We next present a novel technique that delegates this task to the datalog engine itself. The main idea is to extend the datalog program with additional rules that are responsible for the tracking; in this way, the relevant rules and facts can be obtained from the materialisation of the modified program.

Definition 3. *Let $\mathcal{K}$ be a datalog KB and let F be a set of facts in $\text{Mat}(\mathcal{K})$. Then, $\Delta(\mathcal{K}, F)$ is the datalog program containing the rules and facts given next:*

- *each rule and fact in $\mathcal{K}$;*
- *a fact $\bar{P}(\mathbf{a})$ for each fact $P(\mathbf{a})$ in F ;*
- *the following rules for each $r \in \mathcal{K}$ of the form $B_1(\mathbf{x}_1), \dots, B_m(\mathbf{x}_m) \rightarrow H(\mathbf{x})$, and $1 \leq i \leq m$, with $\mathbf{c}_r$ a fresh constant for each r , and $\mathbf{S}$ a fresh predicate:*

$$\bar{H}(\mathbf{x}) \wedge B_1(\mathbf{x}_1) \wedge \dots \wedge B_m(\mathbf{x}_m) \rightarrow \mathbf{S}(\mathbf{c}_r) \quad (1)$$

$$\bar{H}(\mathbf{x}) \wedge B_1(\mathbf{x}_1), \dots \wedge B_m(\mathbf{x}_m) \rightarrow \bar{B}_i(\mathbf{x}_i) \quad (2)$$

The auxiliary predicates $\bar{P}$ are used to record facts involved in proofs; in particular, if $\bar{P}(\mathbf{c})$ is contained in $\text{Mat}(\Delta(\mathcal{K}, F))$, we can conclude that $P(\mathbf{c})$ participates in an SLD proof in $\mathcal{K}$ of a fact in F . Furthermore, each rule $r \in \mathcal{K}$ is represented by a fresh constant $\mathbf{c}_r$, and $\mathbf{S}$ is a fresh predicate that is used to record rules of $\mathcal{K}$ involved in proofs. In particular, if $\mathbf{S}(\mathbf{c}_r)$ is contained in $\text{Mat}(\Delta(\mathcal{K}, F))$, we can conclude that rule r participates in an SLD proof in $\mathcal{K}$ of a fact in F . The additional rules (1) and (2) are responsible for the tracking and make sure that the materialisation of $\Delta(\mathcal{K}, F)$ contains the required information. Indeed, if there is an instantiation $B_1(\mathbf{a}_1) \wedge \dots \wedge B_m(\mathbf{a}_m) \rightarrow H(\mathbf{a})$ of a rule $r \in \Delta$, then, by virtue of (1), $\mathbf{c}_r$ will be added to $\mathbf{S}$, and, by virtue of (2), each $\bar{B}_i(\mathbf{a}_i)$, for $1 \leq i \leq m$, will be derived. Correctness is established as follows.

Theorem 2. *Let $\mathcal{K}$ be a datalog knowledge base and let F be a set of facts in $\text{Mat}(\mathcal{K})$. Then, a fact $P(\mathbf{a})$ (resp. a rule r) in $\mathcal{K}$ participates in an SLD proof of some fact in F iff $\bar{P}(\mathbf{a})$ (resp. $\mathbf{S}(\mathbf{c}_r)$) is in $\text{Mat}(\Delta(\mathcal{K}, F))$.*

7 Summarisation and Answer Dependencies

Once the relevant fragment has been computed, we check, using the fully-fledged reasoner, whether each candidate answer is entailed. This can be computationally expensive if the fragment is large, or there are many candidate answers to verify. To address these issues, we exploit summarisation techniques [4] to efficiently prune candidate answers. The idea behind summarisation is to “shrink” the data by merging constants instantiating the same unary predicates. Since summarisation is equivalent to extending the knowledge base with equality assertions, the summarised knowledge base entails the original one by monotonicity.

Definition 4. *Let $\mathcal{K}$ be a knowledge base. A type T is a set of unary predicates; for a constant a in $\mathcal{K}$, we say that $T = \{A \mid A(a) \in \mathcal{K}\}$ is the type for a . Furthermore, for each type T , let c_T be a globally fresh constant uniquely associated with T . The summary function over $\mathcal{K}$ is the substitution σ mapping each constant a in $\mathcal{K}$ to c_T , where T is the type for a . Finally, the knowledge base $\sigma(\mathcal{K})$ obtained by replacing each constant a in $\mathcal{K}$ with $\sigma(a)$ is called the summary of $\mathcal{K}$.*

By summarising a knowledge base, we overestimate query answers [4].

Proposition 2. *Let $\mathcal{K}$ be a knowledge base, and let σ be the summary function over $\mathcal{K}$. Then, for every query q we have $\sigma(\text{cert}(q, \mathcal{K})) \subseteq \text{cert}(\sigma(q), \sigma(\mathcal{K}))$.*

Summarisation can be exploited to detect spurious answers in G : if a tuple is not in $\text{cert}(\sigma(q), \sigma(\mathcal{K}))$, then it is not in $\text{cert}(q, \mathcal{K})$. Since summarisation can significantly reduce the size of a knowledge base, we can efficiently detect non-answers even if checking them over the summary requires calling the OWL reasoner.

Corollary 1. *Let $\mathcal{K}$ be a knowledge base, let q be a query, let S be a set of tuples, and let $\mathcal{K}' = \mathcal{K}_{[q, S]} \cup \mathcal{K}_\perp$. Furthermore, let σ be the summary function over $\mathcal{K}'$. Then, $\sigma(\mathcal{K}') \not\models \sigma(q(\mathbf{a}))$ implies $\mathcal{K} \not\models q(\mathbf{a})$ for each $\mathbf{a} \in S$.*

Finally, we try to further reduce the calls to the fully-fledged reasoner by exploiting dependencies between the candidate answers. Consider tuples $\mathbf{a}$ and $\mathbf{b}$ in G and the dataset $\mathcal{D}$ in the fragment $\mathcal{K}_{[q, G]} \cup \mathcal{K}_\perp$; furthermore, suppose we can find an endomorphism h of $\mathcal{D}$ in which $h(\mathbf{a}) = \mathbf{b}$. If we can determine (by calling the fully-fledged reasoner) that $\mathbf{b}$ is a spurious answer, then so must be $\mathbf{a}$; as a result, we no longer call the reasoner to check $\mathbf{a}$. We exploit this idea to compute a dependency graph having candidate answers as nodes and an edge $(\mathbf{a}, \mathbf{b})$ whenever an endomorphism in $\mathcal{D}$ exists mapping $\mathbf{a}$ to $\mathbf{b}$. Computing endomorphisms is computationally hard, so we have implemented a sound (but incomplete) greedy algorithm that approximates the dependency graph.

Data	DL	Axioms	Facts
LUBM(n)	<i>SHI</i>	93	$10^5 n$
UOBM ⁻ (n)	<i>SHIN</i>	314	$2 \times 10^5 n$
FLY	<i>SRI</i>	144,407	6,308
DBPedia ⁺	<i>SHOIN</i>	1,757	12,119,662
NPD	<i>SHIF</i>	819	3,817,079

Table 2. Statistics for test data

Strategy	Solved	Univ	$t_{avg.}$
RL Bounds	14	1000	18.4
+ EL Lower Bound	22	1000	11.7
+ Sum, Dep	24	100	29.6

Table 3. Result for LUBM

Strategy	Solved	Univ	$t_{avg.}$
RL Bounds	12	500	0.7
+ Summarisation	14	60	14.0
+ Dependencies	15	1	1.8

Table 4. Result for UOBM⁻

8 Evaluation

We have implemented a prototype system, called PAGOdA, based on RDFS and HermiT (v. 1.3.8). For testing, we used the LUBM and UOBM benchmarks, as well as the Fly Anatomy ontology, DBPedia and NPD FactPages; their key features are summarised in Table 2. Our system, test data, ontologies, and queries are available online.² We compared our system with Pellet (v. 2.3.1) and TrOWL [20] on all datasets. While Pellet is sound and complete, TrOWL relies on approximate reasoning and does not provide correctness guarantees. Tests were performed on a 16 core 3.30GHz Intel Xeon E5-2643 with 125GB of RAM, and running Linux 2.6.32. For each test, we measured materialisation times for upper and lower bound, the time to answer each query, and the number of queries that can be fully answered using different techniques. All times are in seconds.

Materialisation is fast on LUBM [7]: it takes 319s (341s) to materialise the basic lower (upper) bound entailments for LUBM(1000). These bounds match for all 14 standard LUBM queries, and we have used 10 additional queries for which this is not the case; we tested our system on all 24 queries (see Table 3 for a summary of the results). The refined lower bound was materialised in 366s, and it matches the upper bound for 8 of the 10 additional queries; thus, our system could answer 22 of the 24 queries over LUBM(1000) efficiently in 12s on average.³ For the remaining 2 queries, we could scale to LUBM(100) in reasonable time. On LUBM(100) the gaps contain 29 and 14 tuples respectively, none of which were eliminated by summarisation; however, exploiting dependencies between gap tuples reduced the calls to HermiT to only 3 and 1 respectively, with the majority of time taken in extraction (avg. 45s) and HermiT calls (avg. 281s). On LUBM(1000), Pellet ran out of memory. For LUBM(100), Pellet took on average 8.2s to answer the standard queries with an initialisation overhead of 388s. TrOWL timed out after 1h on LUBM(100).

² <http://www.cs.ox.ac.uk/isg/tools/PAGOdA/>

³ Average query answering times are measured after materialisation.

UOBM is an extension of LUBM [25]. Query answering over UOBM requires equality reasoning (e.g., to deal with cardinality constraints), which is not natively supported by RDFS⁴, so we have used a slightly weakened ontology UOBM⁻ for which equality is not required. Materialisation is still fast on UOBM⁻(500): it takes 346s (378s) to materialise the basic lower (upper) bound entailments. We have tested the 15 standard queries (see Table 4). The basic lower and upper bounds match for 12 queries; our system is efficient for these queries, with an average query answering time of less than 1s over UOBM⁻(500). For 2 of the remaining queries, summarisation prunes all candidate answers. Average times for these queries were under 15s for UOBM⁻(60). For the one remaining query, summarisation rules out 6245 among 6509 answers in the gap, and the dependency analysis groups all the remaining individuals. HermiT, however, takes 20s to check the representative answer for UOBM⁻(1), and 4000s for UOBM⁻(10). Pellet times out even on UOBM⁻(1). TrOWL took 237s on average to answer 14 out of the 15 queries over UOBM⁻(60).⁵ Furthermore, a comparison with our system reveals that TrOWL answers may be neither sound nor complete for most test queries.

Fly Anatomy is a complex ontology, rich in existential axioms, and including a dataset with over 6,000 facts. We tested it with five queries provided by the developers the ontology. It took 88s (106s) to materialise lower (upper) bound entailments. The basic lower bounds for all queries are empty, whereas the refined lower bounds (which take 185s to materialise) match with the upper bound in all cases; as a result, we can answer the queries in 0.2s on average. Pellet fails to answer queries given a 1h timeout, and TrOWL returns only empty answers.

In contrast to Fly, the DBPedia dataset is relatively large, but the ontology is simple. To provide a more challenging test, we have used the LogMap ontology matching system [9] to extend DBPedia with the tourism ontology which contains both disjunctive and existential axioms. Since the tested systems report errors on datatypes, we have removed all axioms and facts involving datatypes. It takes 45s (47s) to materialise the basic lower (upper) bound entailments. The upper bound was unsatisfiable and it took 2.6s to check satisfiability of the $\mathcal{K}_\perp$ fragment. We queried for instances of all 441 atomic concepts. Bounds matched in 439 cases (using the refined lower bound), and these queries were answered in 0.3s on average. Summarisation filtered out all gap tuples for the remaining two queries. The answer time for both queries was less than 3s. Pellet takes 280.9s to initialise and answers each query in an average time of 16.2s. TrOWL times out after 1h.

The NPD FactPages ontology describes petroleum activities on the Norwegian continental shelf. The ontology is not Horn, and it includes existential axioms. As in the case of DBPedia, we removed axioms involving datatypes. Its dataset has about 4 million triples; it takes 17s (22s) to materialise the lower (upper) bound entailments. The upper bound is unsatisfiable, and it took 30s to check satisfiability of $\mathcal{K}_\perp$. We queried for the instances of the 329 atomic

⁴ RDFS supports equality via its axiomatisation as a congruence relation.

⁵ An exception is reported for the remaining query.

concepts, and could answer all queries in 2.5s on average. Queries with matching bounds (290 out of 329) could be answered on 0.1s; full reasoning is involved in only 4 queries. Pellet took 127s to initialise, and average query answering time was 3s. TrOWL took 1.3s to answer queries on average; answers were complete for 320 out of the 329 queries.

9 Related Techniques

The SCREECH system [21] exploits the KAON2 reasoner [8] to rewrite a *SHIQ* ontology into disjunctive datalog while preserving atomic queries, and then transforms $\vee$ into $\wedge$; the resulting over-approximation can be used to compute upper bound query answers. This technique is restricted to *SHIQ* ontologies and atomic queries; furthermore, the set of rules obtained from KAON2 can be expensive to compute, as well as of exponential size. Both the QUILL system [17] and the work of [25] under-approximate the ontology into OWL 2 QL; however, neither approximation is independent of both query and data, and using OWL 2 QL increases the chances that the approximated ontology will be unsatisfiable.

The SHER system uses summarisation to efficiently compute an upper bound answer, with exact answers then being computed via successive relaxations [3, 4]. The technique has been shown to be scalable, but it is only known to be applicable to *SHIN* and atomic queries, and is less modular than our approach. In contrast, our approach can profitably exploit the summarisation technique, and could even improve scalability for the hardest queries by replacing Hermit with SHER when the extracted fragment is *SHIN*.

10 Discussion

We have proposed a novel approach for query answering that integrates scalable and complete reasoners to provide pay-as-you-go performance. Our evaluation shows that 768 of the 814 test queries could be answered using highly scalable lower and upper bound computations, 39 of the remaining 46 queries yielded to extraction and summarisation techniques, and even for the remaining 7 queries our fragment extraction and dependency techniques greatly improved scalability. Our approach is complementary to other optimisation efforts, and could immediately benefit from alternative techniques for efficiently computing lower bounds and/or a more efficient OWL reasoner. Our technical results are very general, and hold for any language L captured by generalised rules.

There are still many possibilities for future work. For the immediate future, our main focus will be improving the fragment extraction and checking techniques so as to improve scalability for the hardest queries.

Acknowledgements. This work was supported by the Royal Society, the EPSRC projects Score!, Exoda, and MaSI³, and the FP7 project OPTIQUE.

References

1. Bishop, B., Kiryakov, A., Ognyanoff, D., Peikov, I., Tashev, Z., Velkov, R.: OWLIM: A family of scalable semantic repositories. *Semantic Web* 2(1), 33–42 (2011)
2. Cuenca Grau, B., Motik, B., Stoilos, G., Horrocks, I.: Completeness guarantees for incomplete ontology reasoners: Theory and practice. *J. Artif. Intell. Res. (JAIR)* 43, 419–476 (2012)
3. Dolby, J., Fokoue, A., Kalyanpur, A., Kershenbaum, A., Schonberg, E., Srinivas, K., Ma, L.: Scalable semantic retrieval through summarization and refinement. In: *AAAI*. pp. 299–304 (2007)
4. Dolby, J., Fokoue, A., Kalyanpur, A., Schonberg, E., Srinivas, K.: Scalable highly expressive reasoner (SHER). *J. Web Sem.* 7(4), 357–361 (2009)
5. Eiter, T., Ortiz, M., Simkus, M.: Conjunctive query answering in the description logic SH using knots. *J. Comput. Syst. Sci.* 78(1), 47–85 (2012)
6. Glimm, B., Lutz, C., Horrocks, I., Sattler, U.: Conjunctive query answering for the description logic SHIQ. *J. Artif. Intell. Res. (JAIR)* 31, 157–204 (2008)
7. Guo, Y., Pan, Z., Heflin, J.: LUBM: A benchmark for OWL knowledge base systems. *J. Web Sem.* 3(2-3), 158–182 (2005)
8. Hustadt, U., Motik, B., Sattler, U.: Reasoning in description logics by a reduction to disjunctive datalog. *J. Autom. Reasoning* 39(3), 351–384 (2007)
9. Jiménez-Ruiz, E., Cuenca Grau, B., Zhou, Y., Horrocks, I.: Large-scale interactive ontology matching: Algorithms and implementation. In: *ECAI*. pp. 444–449 (2012)
10. Kollia, I., Glimm, B.: Optimizing SPARQL query answering over OWL ontologies. *J. Artif. Intell. Res. (JAIR)* 48, 253–303 (2013)
11. Kontchakov, R., Lutz, C., Toman, D., Wolter, F., Zakharyashev, M.: The combined approach to ontology-based data access. In: *IJCAI*. pp. 2656–2661 (2011)
12. Möller, R., Neuenstadt, C., Özçep, Ö.L., Wandelt, S.: Advances in accessing big data with expressive ontologies. In: *Description Logics*. pp. 842–853 (2013)
13. Motik, B., Cuenca Grau, B., Horrocks, I., Wu, Z., Fokoue, A., Lutz, C.: OWL 2 Web Ontology Language Profiles. W3C Recommendation (27 October 2009), available at <http://www.w3.org/TR/owl2-profiles/>
14. Motik, B., Nenov, Y., Piro, R., Horrocks, I., Olteanu, D.: Parallel materialisation of datalog programs in main-memory rdf databases. In: *AAAI* (2014)
15. Motik, B., Patel-Schneider, P.F., Parsia, B.: OWL 2 Web Ontology Language Structural Specification and Functional-style Syntax. W3C Recommendation (27 October 2009 2009), available at <http://www.w3.org/TR/owl2-syntax/>
16. Motik, B., Shearer, R., Horrocks, I.: Hypertableau reasoning for description logics. *J. Artif. Intell. Res. (JAIR)* 36, 165–228 (2009)
17. Pan, J.Z., Thomas, E., Zhao, Y.: Completeness guaranteed approximations for OWL-DL query answering. In: *Description Logics* (2009)
18. Sirin, E., Parsia, B., Cuenca Grau, B., Kalyanpur, A., Katz, Y.: Pellet: A practical OWL-DL reasoner. *J. Web Sem.* 5(2), 51–53 (2007)
19. Stefanoni, G., Motik, B., Horrocks, I.: Introducing nominals to the combined query answering approaches for EL. In: *AAAI* (2013)
20. Thomas, E., Pan, J.Z., Ren, Y.: TrOWL: Tractable OWL 2 reasoning infrastructure. In: *ESWC* (2). pp. 431–435 (2010)
21. Tserendorj, T., Rudolph, S., Krötzsch, M., Hitzler, P.: Approximate OWL-reasoning with Screech. In: *RR*. pp. 165–180 (2008)

22. Urbani, J., van Harmelen, F., Schlobach, S., Bal, H.E.: QueryPIE: Backward reasoning for OWL Horst over very large knowledge bases. In: International Semantic Web Conference (1). pp. 730–745 (2011)
23. Urbani, J., Kotoulas, S., Maassen, J., van Harmelen, F., Bal, H.E.: WebPIE: A web-scale parallel inference engine using MapReduce. *J. Web Sem.* 10, 59–75 (2012)
24. W3C SPARQL Working Group: SPARQL 1.1 Overview. W3C Recommendation (21 March 2013), available at <http://www.w3.org/TR/sparql11-overview/>
25. Wandelt, S., Möller, R., Wessel, M.: Towards scalable instance retrieval over ontologies. *Int. J. Software and Informatics* 4(3), 201–218 (2010)
26. Wu, Z., Eadon, G., Das, S., Chong, E.I., Kolovski, V., Annamalai, M., Srinivasan, J.: Implementing an inference engine for RDFS/OWL constructs and user-defined rules in oracle. In: ICDE. pp. 1239–1248 (2008)
27. Zhou, Y., Cuenca Grau, B., Horrocks, I., Wu, Z., Banerjee, J.: Making the most of your triple store: query answering in OWL 2 using an RL reasoner. In: WWW. pp. 1569–1580 (2013)
28. Zhou, Y., Nenov, Y., Cuenca Grau, B., Horrocks, I.: Complete query answering over Horn ontologies using a triple store. In: ISWC (1). pp. 720–736 (2013)
29. Zhou, Y., Nenov, Y., Cuenca Grau, B., Horrocks, I.: Pay-as-you-go OWL query answering using a triple store. In: AAAI (2014)