

Finite Characterizations of $\mathcal{EL}$ Ontologies and Concept Inclusions

Maurice Funk^{1,2}, Simon Hosemann^{1,3} and Carsten Lutz^{1,2,3}

¹Leipzig University

²Center for Scalable Data Analytics and Artificial Intelligence (ScaDS.AI)

³School of Embedded Composite Artificial Intelligence (SECAI)

Abstract

We study the problem of finitely characterizing ontologies, concept inclusions, and concept equivalences, all formulated in the description logic $\mathcal{EL}$, in terms of positive and negative examples. We consider two different kinds of examples: finite interpretations and finite pointed interpretations, with the expressive power of the latter being higher than that of the former. We prove that non-tautological concept inclusions and concept equivalences are characterizable by a finite set of pointed examples, and that acyclic inclusions and equivalences are even characterizable by a finite set of non-pointed examples. For ontologies, in contrast, finite characterizability fails even for pointed examples. This leads us to consider the important class of left-atomic ontologies, where only concept names can appear on the left side of concept inclusions and equivalences. There, finite characterizability in terms of pointed examples is regained.

Keywords

Description Logics, Finite Characterizations, Concept Inclusions, Ontologies, Fitting

1. Introduction

In many synthesis problems, the goal is to construct a logical artifact from a finite set of positive and negative examples [1, 2, 3, 4]. In description logics, this form of synthesis has been studied for concepts [5, 6, 7], ontologies [8, 9, 10], and queries [11, 12], often under the name of ‘fitting problems’. One main motivation for this line of research is to support a knowledge engineer in the construction of a desired artifact. Given that there are typically many non-equivalent artifacts that fit a given set of examples, the question arises how precisely an artifact can actually be described by a finite set of examples. Ideally, it would be possible to describe the artifact uniquely up to equivalence. A finite set of examples that achieves this is called a finite characterization of the artifact.

In this article, we study finite characterizations of ontologies formulated in the description logic $\mathcal{EL}$, as well as of concept inclusions and equivalences, which form their main constituents. In particular, we provide a detailed analysis of when finite characterizations exist and when they do not. As examples, we consider finite interpretations and finite pointed interpretations. In the former, concept inclusions and equivalences are evaluated globally on the whole interpretation whereas in the latter they are evaluated only locally at a distinguished point. While non-pointed interpretations are closer to the usual semantics of ontologies, pointed interpretations have certain advantages. First, in the case of positive examples, they avoid forcing the user to produce very large examples only to achieve global satisfaction of concept inclusions and equivalences. And second, it turns out that the expressive power of pointed examples is higher than that of non-pointed ones. Consequently, finite characterizations in terms of pointed examples exist in several cases where finite characterizations in terms of non-pointed examples are not available.

We point out that finite characterizations may also be of interest for other reasons. For instance, computing a set of examples that characterizes a manually constructed concept inclusion or ontology

 DL 2026: 39th International Workshop on Description Logics, July 17–19, 2026, Lisbon, Portugal

 maurice.funk@uni-leipzig.de (M. Funk); simon.hosemann@uni-leipzig.de (S. Hosemann); carsten.lutz@uni-leipzig.de (C. Lutz)

 0000-0003-1823-9370 (M. Funk); 0009-0004-6243-6114 (S. Hosemann); 0000-0002-8791-6702 (C. Lutz)



© 2026 Copyright for this paper by its authors. Use permitted under Creative Commons License Attribution 4.0 International (CC BY 4.0).

Table 1

Summary of finite characterizability results.

$\mathcal{EL}$ object	By pointed examples	By non-pointed examples
Tautology	no, Prop. 9	no [‡]
Concept inclusion	yes, Thm. 10	open
Concept equivalence	yes, Thm. 13	no, Prop. 14
Acyclic concept equivalence or inclusion	yes [†]	yes, Thm. 15
Ontology	no, Prop. 17	no [‡]
Left-atomic ontology	yes, Thm. 20	no, Prop. 19

[†] By restriction from the preceding row.[‡] From the pointed case via Lemma 4.

can help illustrate to a user the ramifications of their modeling choices. In addition, the existence of finite characterizations is closely connected to active learning in Angluin’s model of exact learnability [13, 14]. In particular, finite characterizability is a necessary condition for exact learnability with membership queries.

Some care is needed regarding the class of artifacts with respect to which a finite characterization is sought. For instance, characterizing a concept inclusion with respect to the class of all concept inclusions is a weaker requirement than characterizing it with respect to the class of all concept equivalences: every concept inclusion $C \sqsubseteq D$ is equivalent to the concept equivalence $C \equiv C \sqcap D$. In the remainder of this article, unless stated otherwise, finite characterizations are always understood to be relative to the same class of artifacts, for example concept inclusions with respect to the class of all concept inclusions.

We begin with concept inclusions and observe that, under the mild assumption that the signature contains at least one role name, tautological concept inclusions are not finitely characterizable, regardless of the form of examples used. The same holds for concept equivalences and ontologies. Apart from tautologies, however, all $\mathcal{EL}$ concept inclusions turn out to be finitely characterizable by finite sets of pointed examples, and the same is true for concept equivalences (and thus also for all concept inclusions w.r.t. all concept equivalences). The characterizations are even in terms of tree interpretations and can be computed in polynomial time. Their construction relies on the notion of the frontier of an $\mathcal{EL}$ concept. In contrast, there are simple concept equivalences that are not finitely characterizable by non-pointed examples. For concept inclusions, finite characterizability by non-pointed examples remains open.

We then turn towards concept inclusions $C \sqsubseteq D$ that are acyclic, meaning that the well-known restricted chase procedure terminates when used with $C \sqsubseteq D$, starting from C viewed as an interpretation. A concept equivalence $C \equiv D$ is acyclic if the inclusions $C \sqsubseteq D$ and $D \sqsubseteq C$ are. We show that in contrast to unrestricted inclusions and equivalences, the acyclic case enjoys finite characterizability in terms of non-pointed examples. Again, the examples are trees and can be computed in polynomial time.

We next consider ontologies. In the unrestricted case, the situation is bleak: even simple ontologies such as $\{A \sqsubseteq \exists r. \top\}$ are not finitely characterizable, regardless of the form of examples used. The apparent contradiction between the non-characterizability of this ontology and the finite characterizability of concept inclusions is explained by the fact that, here, we seek a characterization w.r.t. the class of all ontologies rather than the class of all concept inclusions. We then turn to left-atomic $\mathcal{EL}$ ontologies in which the left side of all concept inclusions and equivalences must be a concept name. Many well-known ontologies, like Gene Ontology and large parts of SNOMED CT are left-atomic.¹ Our main result on ontologies is that left-atomic ontologies are finitely characterizable by pointed examples, w.r.t. the class of all left-atomic ontologies. Again, tree examples suffice, but they are not of polynomial size.

Our results are summarized in Table 1. Full proofs are available in the appendix, provided at <https://www.informatik.uni-leipzig.de/kr/research/papers.html>.

¹In an exploratory check, we analyzed all OWL files linked on the Open Biological and Biomedical Ontology (OBO) Foundry website (<https://obofoundry.org/>). Of 290 distinct files, we successfully retrieved and parsed 243. A naive syntactic check

Related Work

Finite characterizations have been studied for a variety of formalisms, such as conjunctive queries [13], temporal logic formulas [15, 16], schema mappings [17], and modal logic formulas [18].

Closest to our work is recent research on finite characterizability of description logic concepts [19, 14]. Finite characterizability of $\mathcal{EL}$ and $\mathcal{ELI}$ concepts under DL-Lite ontologies by ABox examples is investigated in [19]. A systematic account of the fragments of $\mathcal{ALCQI}$ in which concepts admit finite characterizations by finite interpretations, both without background ontologies and relative to DL-Lite ontologies, is provided in [14]. In contrast to these works, we consider non-pointed examples in addition to pointed ones and study finite characterizations of $\mathcal{EL}$ ontologies, concept inclusions, and concept equivalences, rather than concepts.

Related to finite characterizations are investigations into fitting and learning problems. In this direction, the fitting of $\mathcal{EL}$ and $\mathcal{ELI}$ ontologies as well as several TGD classes to relational structures is studied in [9]. Fitting ontologies to ABox examples, for $\mathcal{ALC}$ and $\mathcal{EL}$, is investigated in [8] and [10]. Another adjacent line of work concerns the exact learnability of $\mathcal{EL}$ ontologies from concept inclusion examples in [20].

2. Preliminaries

Basic Notions. Let N_C and N_R be infinite sets of concept names and role names. $\mathcal{EL}$ *concepts*, or *concepts* for short, are defined by the rule $C, D ::= \top \mid A \mid C \sqcap D \mid \exists r.C$ where $A \in N_C$ and $r \in N_R$. For a finite set of concepts $S = \{C_1, \dots, C_n\}$, we write $\sqcap S$ for $C_1 \sqcap \dots \sqcap C_n$, and define $\sqcap \emptyset := \top$. We use $\text{sub}(C)$ to denote the set of *subconcepts* of a concept C and $\text{rd}(C)$ to denote its *role depth*, that is, the nesting depth of expressions $\exists r.D$ in it. *Concept inclusions* and *concept equivalences* take the form $C \sqsubseteq D$ and $C \equiv D$, respectively, where C and D are concepts. An ontology $\mathcal{O}$ is a finite set of concept inclusions and equivalences. We say that $\mathcal{O}$ is *left-atomic* if the left side of every concept inclusion and equivalence in $\mathcal{O}$ is a concept name.

The semantics of concepts and ontologies is defined as usual in terms of *interpretations* $\mathcal{I} = (\Delta^{\mathcal{I}}, \cdot^{\mathcal{I}})$, where $\Delta^{\mathcal{I}} \neq \emptyset$ is the domain of $\mathcal{I}$ and $\cdot^{\mathcal{I}}$ assigns a set $A^{\mathcal{I}} \subseteq \Delta^{\mathcal{I}}$ to every $A \in N_C$ and a binary relation $r^{\mathcal{I}} \subseteq \Delta^{\mathcal{I}} \times \Delta^{\mathcal{I}}$ to every $r \in N_R$. The *extension* $C^{\mathcal{I}}$ of a concept C is then defined as usual [21]. An interpretation $\mathcal{I}$ *satisfies*, or *is a model of*, a concept inclusion $C \sqsubseteq D$ if $C^{\mathcal{I}} \subseteq D^{\mathcal{I}}$, and a concept equivalence $C \equiv D$ if $C^{\mathcal{I}} = D^{\mathcal{I}}$. We denote this by $\mathcal{I} \models C \sqsubseteq D$ and $\mathcal{I} \models C \equiv D$. Moreover, $\mathcal{I}$ is a *model* of an ontology $\mathcal{O}$ if it satisfies all concept inclusions and equivalences in it, denoted $\mathcal{I} \models \mathcal{O}$.

Let α be a concept inclusion or equivalence, or an ontology. We say that α is a *consequence* of an ontology $\mathcal{O}$, written $\mathcal{O} \models \alpha$, if every model of $\mathcal{O}$ is also a model of α . Ontologies $\mathcal{O}$ and $\mathcal{O}'$ are *equivalent* if $\mathcal{O} \models \mathcal{O}'$ and $\mathcal{O}' \models \mathcal{O}$. For $\mathcal{O} = \emptyset$, we may write $\models \alpha$ in place of $\mathcal{O} \models \alpha$ and we call α a *tautology*. For $\mathcal{O} = \{C \sqsubseteq D\}$, we may write $C \sqsubseteq D \models \alpha$ in place of $\mathcal{O} \models \alpha$.

A *signature* is a set of concept and role names. For any syntactic object α such as a concept inclusion, concept equivalence or an ontology, we use $\text{sig}(\alpha)$ to denote the set of concept and role names used in α .

Fittings and Finite Characterizations. A *pointed interpretation* is a pair $(\mathcal{I}, a)$ where $\mathcal{I}$ is an interpretation and $a \in \Delta^{\mathcal{I}}$. An interpretation is *finite* if $\Delta^{\mathcal{I}}$ is finite and $\cdot^{\mathcal{I}}$ assigns a non-empty extension to only finitely many concept and role names. A pointed interpretation $(\mathcal{I}, a)$ is a *tree* if the underlying directed graph of $\mathcal{I}$ is a directed tree rooted at a . An interpretation $\mathcal{I}$ is a *tree* if $(\mathcal{I}, a)$ is a tree for some $a \in \Delta^{\mathcal{I}}$. A *non-pointed example* is a finite interpretation, and a *pointed example* is a finite pointed interpretation. When it is not important whether an example is pointed or non-pointed, we simply speak of an example. For a pointed example $(\mathcal{I}, a)$ and an ontology $\mathcal{O}$, we write $(\mathcal{I}, a) \models \mathcal{O}$ if $a \in C^{\mathcal{I}}$ implies $a \in D^{\mathcal{I}}$ for all concept inclusions $C \sqsubseteq D \in \mathcal{O}$, and $a \in C^{\mathcal{I}}$ iff $a \in D^{\mathcal{I}}$ for all concept

of ‘subClassOf’ and ‘equivalentClasses’ axioms classified 176 of these 243 files as left-atomic, corresponding to 72% of the analyzed files.

equivalences $C \equiv D \in \mathcal{O}$. A *collection of examples* is a pair (P, N) with P and N finite sets of positive and negative examples, respectively. An ontology $\mathcal{O}$ *fits* P if $e \models \mathcal{O}$ for all $e \in P$, it *fits* N if $e \not\models \mathcal{O}$ for all $e \in N$, and it *fits* (P, N) if it fits both P and N .

Example 1. In the case of pointed examples, fitting is not invariant under logical equivalence. Consider the concept inclusion $A \sqsubseteq \exists r.A$ and the collection of pointed examples $(\{(\mathcal{I}, a_1)\}, \emptyset)$ with $\mathcal{I}$ as depicted below.

$$A \quad a_1 \longrightarrow a_2 \quad A$$

While the concept inclusion $A \sqsubseteq \exists r.(A \sqcap \exists r.A)$ is equivalent to $A \sqsubseteq \exists r.A$, it does not fit $(\{(\mathcal{I}, a_1)\}, \emptyset)$, since $(\mathcal{I}, a_1) \not\models \{A \sqsubseteq \exists r.(A \sqcap \exists r.A)\}$.

Also observe that if an ontology $\mathcal{O}$ fits (P, N) , then this does not imply that every $\alpha \in \mathcal{O}$ fits (P, N) . In fact, $\mathcal{O}$ fits (P, N) if and only if for every $e \in N$, there is an $\alpha \in \mathcal{O}$ such that α fits $(P, \{e\})$.

Definition 2 (Finite Characterizability). Let $\mathcal{C}$ be a class of ontologies. A collection (P, N) of examples is a *characterization* of an ontology $\mathcal{O}$ w.r.t. $\mathcal{C}$ if $\mathcal{O}$ fits (P, N) and every ontology $\mathcal{O}' \in \mathcal{C}$ with $\text{sig}(\mathcal{O}') \subseteq \text{sig}(\mathcal{O})$ that fits (P, N) is equivalent to $\mathcal{O}$. An ontology $\mathcal{O}$ is *finitely characterizable* w.r.t. $\mathcal{C}$ by non-pointed examples if there exists a finite collection of non-pointed examples (P, N) that is a characterization of $\mathcal{O}$ w.r.t. $\mathcal{C}$, and likewise for finite characterizability by pointed examples.

The notions introduced in Definition 2 also apply to concept inclusions and equivalences when these are viewed as singleton ontologies. When we say that every ontology from a class of ontologies $\mathcal{C}$ is finitely characterizable w.r.t. $\mathcal{C}$, we usually omit the ‘w.r.t. $\mathcal{C}$ ’ qualifier, and likewise for concept inclusions and equivalences. We remark that the condition $\text{sig}(\mathcal{O}') \subseteq \text{sig}(\mathcal{O})$ in Definition 2 aligns with many ontology learning scenarios where the signature of the target ontology is assumed to be known and fixed in advance. This restriction also makes finite characterizations more likely to exist, since otherwise one can often construct non-equivalent fitting ontologies that use additional symbols.

Example 3. Consider the left-atomic ontology $\mathcal{O} = \{A \sqsubseteq \exists r.A\}$ and the collection of pointed examples (P, N) , where $P = \{(\mathcal{I}, a_1), (\mathcal{J}, b_1)\}$ and $N = \{(\mathcal{K}, c_1)\}$ with $\mathcal{I}, \mathcal{J}$, and $\mathcal{K}$ as depicted below.

$$\begin{array}{ccc} \mathcal{I} & \begin{array}{c} a_1 \quad A \\ \downarrow \\ a_2 \quad A \end{array} & \mathcal{J} \quad \begin{array}{c} b_1 \\ \downarrow \\ b_2 \quad A \end{array} & \mathcal{K} \quad \begin{array}{c} c_1 \quad A \\ \downarrow \\ c_2 \end{array} \end{array}$$

Then (P, N) is a finite characterization of $\mathcal{O}$ w.r.t. left-atomic ontologies. Observe that this no longer holds when the positive example $(\mathcal{I}, a_1)$ is dropped, since then $\{A \sqsubseteq \exists r.\exists r.\top\}$ also fits. It also no longer holds when the negative example is dropped, since then $\{A \sqsubseteq \exists r.\top\}$ also fits.

Now consider the left-atomic ontology $\mathcal{O}' = \mathcal{O} \cup \{B \equiv \exists r.\exists r.\top\}$. While $\mathcal{O}'$ also fits (P, N) and $\mathcal{O} \not\models \mathcal{O}'$, we have $\text{sig}(\mathcal{O}') \not\subseteq \text{sig}(\mathcal{O})$. This illustrates the relevance of the restriction to $\text{sig}(\mathcal{O})$ in Definition 2.

We next observe that the expressive power of pointed examples is at least as high as that of non-pointed examples.

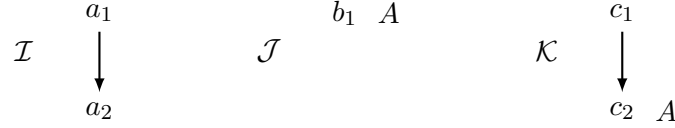
Lemma 4. If an ontology $\mathcal{O}$ is finitely characterizable w.r.t. a class of ontologies $\mathcal{C}$ by a collection of non-pointed examples (P, N) , then $\mathcal{O}$ is finitely characterizable w.r.t. $\mathcal{C}$ by a collection of pointed examples (P', N') . Moreover, given (P, N) and $\mathcal{O}$ we can compute (P', N') in polynomial time.

Proof. Let (P, N) be a finite collection of non-pointed examples that characterizes $\mathcal{O}$ w.r.t. $\mathcal{C}$. Set $P' = \bigcup_{\mathcal{I} \in P} \{(\mathcal{I}, a) \mid a \in \Delta^{\mathcal{I}}\}$ and $N' = \bigcup_{\mathcal{J} \in N} \{(\mathcal{J}, b) \mid b \in \Delta^{\mathcal{J}}, (\mathcal{J}, b) \not\models \mathcal{O}\}$. Clearly, (P', N') can be constructed in polynomial time and is a collection of pointed examples. From the definition of (P', N') it is straightforward to see that $\mathcal{O}$ fits (P', N') . To show that (P', N') characterizes $\mathcal{O}$ w.r.t. $\mathcal{C}$, we need to prove that any ontology from $\mathcal{C}$ that fits (P', N') is equivalent to $\mathcal{O}$.

Let $\mathcal{O}' \in \mathbf{C}$ be an ontology that fits (P', N') . We show that $\mathcal{O}'$ also fits (P, N) , which entails that $\mathcal{O}'$ is equivalent to $\mathcal{O}$ since (P, N) characterizes $\mathcal{O}$ w.r.t. $\mathbf{C}$. Let $\mathcal{I} \in P$. As $\mathcal{O}'$ fits P' and, for every $a \in \Delta^{\mathcal{I}}$, $(\mathcal{I}, a) \in P'$, $(\mathcal{I}, a) \models \mathcal{O}'$ and therefore $\mathcal{I} \models \mathcal{O}'$. Now let $\mathcal{J} \in N$. Then there must be some $b \in \Delta^{\mathcal{J}}$ with $(\mathcal{J}, b) \in N'$. Since $\mathcal{O}'$ fits N' , we must have $(\mathcal{J}, b) \not\models \mathcal{O}'$. This implies $\mathcal{J} \not\models \mathcal{O}'$, as required. $\square$

The converse of Lemma 4 does not hold, as illustrated by Example 5.

Example 5. Consider the concept equivalence $A \equiv \exists r.A$ and the interpretations $\mathcal{I}$, $\mathcal{J}$, and $\mathcal{K}$ as depicted below.



Then (P, N) with $P = \{(\mathcal{I}, a_1)\}$ and $N = \{(\mathcal{J}, b_1), (\mathcal{K}, c_1)\}$ is a finite characterization of $A \equiv \exists r.A$. (P, N) is obtained from the construction used in the proof of Theorem 13, with redundant examples omitted for simplicity. Now consider the collection of non-pointed examples $(P', N') = (\{\mathcal{I}\}, \{\mathcal{J}, \mathcal{K}\})$. While $A \equiv \exists r.A$ fits (P', N') , the concept equivalence $A \equiv \exists r.\exists r.\top$ also does. This is not accidental. Proposition 14 below shows that $A \equiv \exists r.A$ does not admit a finite characterization by non-pointed examples.

Frontiers. The subsequent constructions of finite characterizations by pointed examples rely on the notions of canonical pointed tree models and frontiers of $\mathcal{EL}$ concepts, which we recall next. For an $\mathcal{EL}$ concept C , let $(\mathcal{I}_C, a_C)$ denote the *canonical pointed tree model* of C , that is, the (finite) pointed tree interpretation that can be constructed in a straightforward recursive way from the structure of C and is rooted at a_C [22]. We shall use the following well-known property of such models.

Lemma 6. For all $\mathcal{EL}$ concepts C, D , $a_C \in D^{\mathcal{I}_C}$ if and only if $\models C \sqsubseteq D$.

Intuitively, the frontier of a concept C captures all possible ways to weaken it.

Definition 7 (Frontier). A frontier of an $\mathcal{EL}$ concept C is a set of $\mathcal{EL}$ concepts $\mathcal{F}$ such that

1. $\models C \sqsubseteq D$ and $\not\models D \sqsubseteq C$ for all $D \in \mathcal{F}$,
2. for all $\mathcal{EL}$ concepts E , if $\models C \sqsubseteq E$ and $\not\models E \sqsubseteq C$, then there is a $D \in \mathcal{F}$ with $\models D \sqsubseteq E$.

For a concept $C = C_1 \sqcap \dots \sqcap C_n$, where each C_i is of the form A , $\top$, or $\exists r.D$, we define $\text{conj}(C) := \{C_i \mid i \in [n], C_i \neq \top\}$. We call C *minimal* if there are no distinct $i, j \in [n]$ with $\models C_i \sqsubseteq C_j$ and all proper subconcepts of C are minimal. It is well-known and easy to see that every $\mathcal{EL}$ concept is equivalent to a unique minimal $\mathcal{EL}$ concept that can be computed in polynomial time [23].

Frontiers of $\mathcal{EL}$ concepts always exist and can be computed in a straightforward way if the concept is minimal. For all minimal concepts C , inductively define a function $\text{Frontier}(C)$ as follows:

$$\begin{aligned}
 \text{Frontier}(\top) &= \emptyset, \\
 \text{Frontier}(A) &= \{\top\}, \\
 \text{Frontier}(C_1 \sqcap C_2) &= \{C_1 \sqcap D \mid D \in \text{Frontier}(C_2)\} \cup \{D \sqcap C_2 \mid D \in \text{Frontier}(C_1)\}, \\
 \text{Frontier}(\exists r.C) &= \{\sqcap\{\exists r.D \mid D \in \text{Frontier}(C)\}\}.
 \end{aligned}$$

Lemma 8 ([23]). For all minimal $\mathcal{EL}$ concepts C , $\text{Frontier}(C)$ is a frontier of C and can be computed in polynomial time.

We write $\text{Frontier}(C)$ also for non-minimal concepts C , meaning $\text{Frontier}(C')$ for the unique minimal $\mathcal{EL}$ concept C' that is equivalent to C .

For $\mathcal{EL}$ concepts C , frontiers are intimately related to finite characterizations. Consider the sets

$$P = \{(\mathcal{I}_C, a_C)\} \quad \text{and} \quad N = \{(\mathcal{I}_E, a_E) \mid E \in \text{Frontier}(C)\}.$$

Using the properties of frontiers (Definition 7) and of canonical pointed tree models, it is easy to verify that (P, N) constitutes a finite characterization of C , in the sense that every concept D that fits (P, N) is equivalent to C : D fitting P implies $\models C \sqsubseteq D$ and, if $\not\models D \sqsubseteq C$, then some $E \in \text{Frontier}(C)$ satisfies $\models E \sqsubseteq D$, contradicting that D fits N .

3. Concept Inclusions and Equivalences

We investigate finite characterizations of concept inclusions and equivalences. We begin with tautologies and show that, as soon as the signature contains a role name, finite characterizability fails. We then turn to the non-tautological case, showing that then both concept inclusions and equivalences are finitely characterizable by pointed examples. In contrast, concept equivalences are not finitely characterizable by non-pointed examples, and the analogous question for concept inclusions remains open. This leads us to the natural subclass of acyclic concept inclusions and equivalences, for which we obtain finite characterizability even in the non-pointed setting.

Proposition 9. *No tautological $\mathcal{EL}$ concept inclusion that contains at least one role name is finitely characterizable, neither by pointed nor by non-pointed examples.*

Proof. Let $C \sqsubseteq D$ be a tautological concept inclusion and let $r \in \text{sig}(C \sqsubseteq D) \cap \mathbf{N}_R$. By Lemma 4 it suffices to show that $C \sqsubseteq D$ is not finitely characterizable w.r.t. all concept inclusions by pointed examples. Assume towards a contradiction that the collection of pointed examples (P, N) is a finite characterization of $C \sqsubseteq D$ w.r.t. all concept inclusions. Since $C \sqsubseteq D$ is a tautology, we must have $N = \emptyset$.

Let $S = \{(\mathcal{I}, a) \in P \mid \exists n \in \mathbb{N}: a \notin (\exists r^n. \top)^\mathcal{I}\}$. If $S = \emptyset$, let $n = 0$. Otherwise, choose n maximal such that for some $(\mathcal{I}, a) \in S$, the element a has an outgoing r -path of length n . Then, by maximality of n , every $(\mathcal{I}, a) \in S$ satisfies $a \notin (\exists r^{n+1}. \top)^\mathcal{I}$. Moreover, every $(\mathcal{I}, a) \in P \setminus S$ satisfies $a \in (\exists r^{n+2}. \top)^\mathcal{I}$.

Hence the concept inclusion $\exists r^{n+1}. \top \sqsubseteq \exists r^{n+2}. \top$ fits (P, N) . Since this concept inclusion is not a tautology, it is not equivalent to $C \sqsubseteq D$, contradicting the assumption that $C \sqsubseteq D$ is finitely characterizable. $\square$

Since every concept inclusion $C \sqsubseteq D$ is equivalent to the concept equivalence $C \equiv C \sqcap D$ and can be viewed as an ontology $\{C \sqsubseteq D\}$, Proposition 9 also implies that no tautological $\mathcal{EL}$ concept equivalence that contains at least one role name is finitely characterizable w.r.t. all $\mathcal{EL}$ concept inclusions, as well as an analogous result for tautological $\mathcal{EL}$ ontologies. In light of Example 1 we remark that $C \sqsubseteq D$ and $C \equiv C \sqcap D$ are even equivalent in the stronger sense that for every pointed example e , we have $e \models C \sqsubseteq D$ if and only if $e \models C \equiv C \sqcap D$.

The main result of this section is that concept equivalences can be finitely characterized by pointed examples w.r.t. all concept equivalences, which implies the same result for concept inclusions. Since, however, concept inclusions are more basic than concept equivalences and more widely used, we first showcase the construction of a characterizing collection of examples for concept inclusions. The following is our main result regarding concept inclusions.

Theorem 10. *Every non-tautological $\mathcal{EL}$ concept inclusion is finitely characterizable by pointed tree examples. Moreover, such a finite characterization can be computed in polynomial time.*

As for the proof of Theorem 10, let $C \sqsubseteq D$ be a non-tautological concept inclusion. Define

$$P = \{(\mathcal{I}_{C \sqcap D}, a_{C \sqcap D})\} \cup \{(\mathcal{I}_E, a_E) \mid E \in \text{Frontier}(C)\} \text{ and} \\ N = \{(\mathcal{I}_C, a_C)\} \cup \{(\mathcal{I}_E, a_E) \mid E \in \text{Frontier}(C \sqcap D) \text{ and } \models E \sqsubseteq C\}.$$

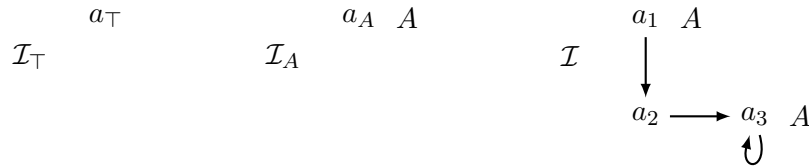
By Lemma 8, the collection (P, N) can clearly be computed in polynomial time. Up to minor details, the above construction is the specialization of the more general construction for concept equivalences given in the proof of Theorem 13 below. Nevertheless, we prove in the appendix that the above construction indeed yields a finite characterization. Here we provide some intuitions. The first negative example ensures that all concept inclusions $C' \sqsubseteq D'$ that fit (P, N) satisfy $\models C \sqsubseteq C'$ and are not tautologies. The second kind of positive example then gives $\models C' \equiv C$. The first positive example ensures that $\models C \sqcap D \sqsubseteq D'$ and the second kind of negative example ensures that $\models C \sqcap D' \sqsubseteq D$. This then implies that $C \sqsubseteq D$ and $C' \sqsubseteq D'$ are equivalent.

Example 11. Consider the concept inclusion $A \sqsubseteq \exists r.A$. The finite characterization of $A \sqsubseteq \exists r.A$ by pointed tree examples following the above construction is (P, N) with

$$P = \{(\mathcal{I}_{A \sqcap \exists r.A}, a_{A \sqcap \exists r.A}), (\mathcal{I}_\top, a_\top)\} \quad \text{and} \quad N = \{(\mathcal{I}_A, a_A), (\mathcal{I}_{A \sqcap \exists r.\top}, a_{A \sqcap \exists r.\top})\}.$$

We leave it as an open problem whether every non-tautological $\mathcal{EL}$ concept inclusion is characterizable by a finite collection of non-pointed examples. We give an example showing that some cyclic concept inclusions admit finite characterizations by non-pointed non-tree examples.

Example 12. Consider again $A \sqsubseteq \exists r.A$ and the interpretations $\mathcal{I}_\top$, $\mathcal{I}_A$, and $\mathcal{I}$ as depicted below.



Then (P, N) with $P = \{\mathcal{I}_\top\}$ and $N = \{\mathcal{I}_A, \mathcal{I}\}$ is a finite characterization of $A \sqsubseteq \exists r.A$. Similarly to the construction in the pointed case, $\mathcal{I}_\top$ and $\mathcal{I}_A$ control the left side of any fitting concept inclusion. $\mathcal{I}_A \in N$ leaves as possible left sides only A and $\top$, and $\mathcal{I}_\top \in P$ excludes $\top$. So the left side of any concept inclusion fitting (P, N) is equivalent to A . The negative example $\mathcal{I}$ then excludes all concept inclusions $A \sqsubseteq C$ with $A \sqsubseteq \exists r.A \models A \sqsubseteq C$ but $A \sqsubseteq C \not\models A \sqsubseteq \exists r.A$. This is the case as any concept inclusion $A \sqsubseteq C$ is satisfied at a_2 and a_3 and violated at a_1 if and only if $\models C \sqsubseteq \exists r.A$. Observe that $\mathcal{I}$ thus excludes an infinite number of concept inclusions that are strictly weaker than $A \sqsubseteq \exists r.A$.

What follows is the main result of this section concerning pointed examples.

Theorem 13. Every non-tautological $\mathcal{EL}$ concept equivalence is finitely characterizable by pointed tree examples. Moreover, such a finite characterization can be computed in polynomial time.

Theorem 13 is proved as follows. Let $C \equiv D$ be a non-tautological concept equivalence. We distinguish three cases:

1. Both $\not\models C \sqsubseteq D$ and $\not\models D \sqsubseteq C$ hold. Then set

$$P = \{(\mathcal{I}_E, a_E) \mid E \in \text{Frontier}(C) \cup \text{Frontier}(D)\} \quad \text{and} \quad N = \{(\mathcal{I}_C, a_C), (\mathcal{I}_D, a_D)\}.$$

2. $\not\models C \sqsubseteq D$ and $\models D \sqsubseteq C$. Then set

$$P = \{(\mathcal{I}_{C \sqcap D}, a_{C \sqcap D})\} \cup \{(\mathcal{I}_E, a_E) \mid E \in \text{Frontier}(C)\} \cup \\ \{(\mathcal{I}_E, a_E) \mid E \in \text{Frontier}(D) \text{ and } \not\models E \sqsubseteq C\}, \\ N = \{(\mathcal{I}_C, a_C)\} \cup \{(\mathcal{I}_E, a_E) \mid E \in \text{Frontier}(D) \text{ and } \models E \sqsubseteq C\}.$$

3. $\models C \sqsubseteq D$ and $\not\models D \sqsubseteq C$. Then P and N are defined as in Case 2, but with C and D swapped.

We compare the above with the construction for concept inclusions. To do this, a concept inclusion $C \sqsubseteq D$ must be viewed as the concept equivalence $C \equiv C \sqcap D$, and thus all concepts D in the above construction must actually be read as $C \sqcap D$. If $C \sqsubseteq D$ is not a tautology, we are therefore clearly in Case 2 above. Moreover, the negative examples of the two constructions are essentially identical. Note that since D must be read as $C \sqcap D$, we have $E \in \text{Frontier}(C \sqcap D)$ with the same qualifier $\models E \sqsubseteq C$. Also note that there is an additional third kind of positive example in the above construction. This is necessary because here we work w.r.t. the class of all concept equivalences, even if we aim to characterize a concept inclusion, whereas for Theorem 10 we work w.r.t. the class of concept inclusions only.

Theorem 13 establishes finite characterizability of $\mathcal{EL}$ concept equivalences by pointed examples. In contrast, finite characterizability by non-pointed examples fails.

Proposition 14. *The concept equivalence $A \equiv \exists r.A$ is not finitely characterizable by non-pointed examples.*

Proof. Assume towards a contradiction that there is a collection of non-pointed examples (P, N) that characterizes $A \equiv \exists r.A$. Consider the concept equivalences $A \equiv \exists r^i.A$ for all $i \geq 2$. As $A \equiv \exists r.A \models A \equiv \exists r^i.A$ for every $i \geq 2$, all these concept equivalences also fit the positive examples in P . But as $A \equiv \exists r^i.A \not\models A \equiv \exists r.A$ for any $i \geq 2$, there must be, for every $i \geq 2$, a negative example $\mathcal{I} \in N$ such that $\mathcal{I} \models A \equiv \exists r^i.A$. Since N is finite and there are infinitely many prime numbers, there must in fact be a single $\mathcal{I} \in N$ such that

(*) $\mathcal{I} \models A \equiv \exists r^p.A$ for infinitely many prime numbers p .

To obtain a contradiction, it suffices to show that this is not possible.

Let $\mathcal{I}$ be a negative example that satisfies (*). Since $\mathcal{I}$ is a negative example, we have $\mathcal{I} \not\models A \equiv \exists r.A$. Choose two distinct primes p_1, p_2 such that $\mathcal{I} \models A \equiv \exists r^{p_i}.A$ for $i \in \{1, 2\}$. For any element $a \in \Delta^{\mathcal{I}}$, define

$$D_a = \{i \in \mathbb{N} \mid a \in (\exists r^i.A)^{\mathcal{I}}\}.$$

Now, as $\mathcal{I} \not\models A \equiv \exists r.A$, there must be an element $a \in \Delta^{\mathcal{I}}$ such that

1. $a \in A^{\mathcal{I}}$ and $a \notin (\exists r.A)^{\mathcal{I}}$ or
2. $a \notin A^{\mathcal{I}}$ and $a \in (\exists r.A)^{\mathcal{I}}$.

In Case 1, it must be that $k_1 p_1 + k_2 p_2 \in D_a$ for all $k_1, k_2 \geq 0$. As p_1 and p_2 are coprime, the Frobenius number exists [24, 25]. It is $p_1 p_2 - p_1 - p_2$ and satisfies the property that $i \in D_a$ for all $i > p_1 p_2 - p_1 - p_2$. Thus, for every $i > p_1 p_2 - p_1 - p_2$, there must be an r -successor b of a such that $i - 1 \in D_b$. However, since $a \notin (\exists r.A)^{\mathcal{I}}$, it follows that $b \notin A^{\mathcal{I}}$. Thus, $\mathcal{I} \not\models A \equiv \exists r^{(i-1)}.A$ for all $i > p_1 p_2 - p_1 - p_2$. This contradicts the fact that $\mathcal{I}$ satisfies (*).

In Case 2, there must be an r -successor b of a such that $b \in A^{\mathcal{I}}$. Then we can argue similarly to Case 1 that $i \in D_b$ for every $i > p_1 p_2 - p_1 - p_2$. Thus, $i + 1 \in D_a$ for every $i > p_1 p_2 - p_1 - p_2$, and hence $\mathcal{I} \models A \equiv \exists r^{(i+1)}.A$ for all $i > p_1 p_2 - p_1 - p_2$. This is again a contradiction to $\mathcal{I}$ satisfying (*).

Note that at no point is the finiteness of $\mathcal{I}$ assumed. Hence the same proof also shows that $A \equiv \exists r.A$ is not finitely characterizable by finite sets of potentially infinite interpretations. $\square$

Central to the proof of Proposition 14 is the cyclic behavior of $A \equiv \exists r.A$. We now consider certain natural and rather general classes of *acyclic* concept inclusions and equivalences and show that these admit finite characterizations in terms of non-pointed examples, even w.r.t. all concept equivalences.

We say that a concept inclusion $C \sqsubseteq D$ is *acyclic* if there is no concept $\exists r.D' \in \text{sub}(D)$ such that $\models D' \sqsubseteq C$. Informally, this is equivalent to demanding that the well-known restricted chase procedure terminates when applied to $\mathcal{I}_C$ under $C \sqsubseteq D$. A concept equivalence $C \equiv D$ is *acyclic* if the concept

inclusions $C \sqsubseteq D$ and $D \sqsubseteq C$ are both acyclic. It is not hard to see that if a concept inclusion $C \sqsubseteq D$ is acyclic, then the equivalent concept equivalence $C \equiv C \sqcap D$ is also acyclic. In particular, there is no subconcept $\exists r.E \in \text{sub}(C \sqcap D)$ with $\models E \sqsubseteq C$: we would have $E \in \text{sub}(C)$ or $E \in \text{sub}(D)$ and the former can easily be shown impossible using an argument based on role depth, while the latter is impossible because $C \sqsubseteq D$ is acyclic. Conversely, there is no subconcept $\exists r.E \in \text{sub}(C)$ with $\models E \sqsubseteq C \sqcap D$ for reasons of role depth.

Theorem 15. *Every non-tautological acyclic $\mathcal{EL}$ concept equivalence is finitely characterizable by non-pointed tree examples w.r.t. all $\mathcal{EL}$ concept equivalences. Moreover, such a finite characterization can be computed in polynomial time.*

Theorem 15 also yields an analogous result for concept inclusions because every acyclic concept inclusion is equivalent to an acyclic concept equivalence.

Theorem 15 is proved as follows. Let $C \equiv D$ be a non-tautological acyclic concept equivalence and let (P, N) be the finite collection of pointed examples that characterizes $C \equiv D$ obtained from the construction in the proof of Theorem 13. Consider the collection of non-pointed examples (P', N') with

$$P' = \{\mathcal{I} \mid (\mathcal{I}, a) \in P\} \cup \{\mathcal{I}_{C \sqcap D}\} \quad \text{and} \quad N' = \{\mathcal{J} \mid (\mathcal{J}, b) \in N\}.$$

It is clear that $C \equiv D$ fits N' . As it turns out, the same is true for P' since $C \equiv D$ is acyclic.²

Lemma 16. *$C \equiv D$ fits P' .*

Proof. We show that $\mathcal{I}_{C \sqcap D} \models C \equiv D$. Assume for contradiction that there is an $a_E \in \Delta^{\mathcal{I}_{C \sqcap D}}$ with $a_E \in C^{\mathcal{I}_{C \sqcap D}}$ and $a_E \notin D^{\mathcal{I}_{C \sqcap D}}$. The case with $a_E \notin C^{\mathcal{I}_{C \sqcap D}}$ and $a_E \in D^{\mathcal{I}_{C \sqcap D}}$ is symmetric. Clearly E cannot be $C \sqcap D$. If E is a true subconcept of C , then a_E cannot be in $C^{\mathcal{I}_{C \sqcap D}}$ for reasons of role-depth. If E is a true subconcept of D , then there is a true subconcept E of D such that $\models E \sqsubseteq C$, contradicting acyclicity. Similar arguments apply to the other kinds of positive examples, additionally using the properties of frontiers and can be found in the proof of Theorem 15 in the appendix. $\square$

Since (P, N) characterizes $C \equiv D$ and $C \equiv D$ fits (P', N') , it should be clear that we can show (P', N') to be a characterization of $C \equiv D$ by proving that every concept equivalence that fits (P', N') also fits (P, N) . This is done in the appendix, and it relies on the additional positive example in P' .

4. Ontologies

We move on to investigate the finite characterizability of $\mathcal{EL}$ ontologies, which generalizes the cases of concept inclusions and concept equivalences considered before. Thus the negative result of Proposition 14 applies also here, that is, $\mathcal{EL}$ ontologies cannot be finitely characterized by non-pointed examples. In contrast to concept equivalences, however, they also cannot be characterized by pointed examples. Even worse, not even acyclic concept inclusions can be characterized w.r.t. the class of $\mathcal{EL}$ ontologies, in the sense made precise by the following proposition.

Proposition 17. *The ontology $\{A \sqsubseteq \exists r.\top\}$ is not finitely characterizable, neither by pointed nor by non-pointed examples.*

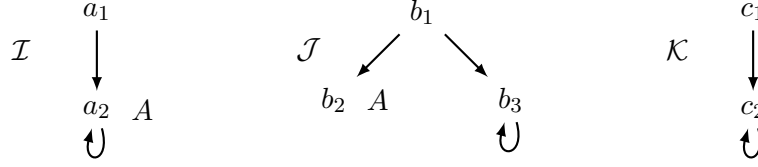
Proof. By Lemma 4 it suffices to prove that $\{A \sqsubseteq \exists r.\top\}$ is not finitely characterizable by pointed examples. Let $\mathcal{O} = \{A \sqsubseteq \exists r.\top\}$. Assume, contrary to what we aim to show, that (P, N) is a pointed collection that finitely characterizes $\mathcal{O}$.

For every example $(\mathcal{I}, a) \in P$, either the length of outgoing r -paths starting in a is bounded by some $n \geq 0$, or $(\mathcal{I}, a)$ satisfies $a \in (\exists r^n.\top)^{\mathcal{I}}$ for all $n \in \mathbb{N}$. Let $n \in \mathbb{N}$ be minimal such that all $(\mathcal{I}, a) \in P$ satisfy $a \in (\exists r^{n+1}.\top)^{\mathcal{I}}$ or $a \notin (\exists r^n.\top)^{\mathcal{I}}$. Then $\mathcal{O}' = \mathcal{O} \cup \{\exists r^n.\top \sqsubseteq \exists r^{n+1}.\top\}$ also fits (P, N) but clearly $\mathcal{O} \not\models \mathcal{O}'$, contradicting the assumption that $\mathcal{O}$ is finitely characterizable. $\square$

²This clearly does not hold if $C \equiv D$ fails to be acyclic because all positive examples in P are finite trees.

We also give an example of a (very simple) ontology that admits a finite characterization in terms of pointed examples. In contrast to all other positive cases in this paper, it requires non-tree examples in the characterization.

Example 18. Consider $\mathcal{O} = \{\top \sqsubseteq \exists r.A\}$ and the collection of pointed examples (P, N) with $P = \{(\mathcal{I}, a_1), (\mathcal{J}, b_1)\}$, $N = \{(\mathcal{I}_\top, a_\top), (\mathcal{K}, c_1)\}$ and the interpretations $\mathcal{I}, \mathcal{J}, \mathcal{K}$ as depicted below.



We show in the appendix that (P, N) is a finite characterization of $\mathcal{O}$.

We next consider the class of left-atomic $\mathcal{EL}$ ontologies that play an important role in practice and, as it turns out, can be finitely characterized.

Proposition 17 shows that left-atomic ontologies are not finitely characterizable w.r.t. the class of all $\mathcal{EL}$ ontologies, neither by pointed nor by non-pointed examples. In the non-pointed case, this can be strengthened to be w.r.t. the class of all left-atomic ontologies. This in fact follows from the proof of Proposition 14 which establishes non-characterizability of the concept equivalence $A \equiv \exists r.A$ not only w.r.t. the class of all concept equivalences, but even w.r.t. the class of all left-atomic concept equivalences.

Proposition 19. The left-atomic ontology $\{A \equiv \exists r.A\}$ is not finitely characterizable by non-pointed examples.

Moving to pointed examples improves the situation. In the following, we consider the finite characterizability of left-atomic ontologies w.r.t. the class of all left-atomic ontologies and show the following main result of this section.

Theorem 20. Every left-atomic $\mathcal{EL}$ ontology is finitely characterizable by pointed tree examples.

To prepare for the proof of Theorem 20, we extend the notion of role depth to concept inclusions, concept equivalences and ontologies in the natural way: $\text{rd}(C \sqsubseteq D) = \text{rd}(C \equiv D) = \max(\text{rd}(C), \text{rd}(D))$ and $\text{rd}(\mathcal{O}) = \max\{\text{rd}(\alpha) \mid \alpha \in \mathcal{O}\}$. In what follows, we assume without loss of generality that left-atomic ontologies $\mathcal{O}$ contain only concept equivalences and that for every $A \in \text{sig}(\mathcal{O})$, $\mathcal{O}$ contains at least one concept equivalence $A \equiv C \in \mathcal{O}$. This is possible because every left-atomic concept inclusion $A \sqsubseteq C$ is equivalent to the concept equivalence $A \equiv A \sqcap C$ and the concept equivalence $A \equiv A$ is a tautology.

The central property of left-atomic ontologies that we use to obtain a finite characterization by pointed examples is the following.

Lemma 21. For all left-atomic $\mathcal{EL}$ ontologies $\mathcal{O}, \mathcal{O}'$ such that $\text{sig}(\mathcal{O}') \cap \text{N}_C \subseteq \text{sig}(\mathcal{O})$ and $\mathcal{O} \not\models \mathcal{O}'$, there exists an $\mathcal{EL}$ concept C such that $\text{sig}(C) \subseteq \text{sig}(\mathcal{O})$, $\text{rd}(C) \leq \text{rd}(\mathcal{O})$, $(\mathcal{I}_C, a_C) \models \mathcal{O}$, and $(\mathcal{I}_C, a_C) \not\models \mathcal{O}'$.

The main virtue of Lemma 21 is that the role depth of C is bounded by that of $\mathcal{O}$, not of $\mathcal{O}'$. To see that this is helpful for a finite characterization, think of $\mathcal{O}$ as the ontology that we want to characterize and of $\mathcal{O}'$ as ranging over infinitely many non-equivalent ontologies that we want to exclude. Also note that Lemma 21 is specific to left-atomic ontologies. For ontologies $\mathcal{O}, \mathcal{O}'$ that are not required to be left-atomic, we also find a concept C as in the lemma, but its role depth depends on that of $\mathcal{O}'$. This is because we want to falsify $\mathcal{O}'$ and thus need to make the left-hand side of some concept inclusion or equivalence in $\mathcal{O}'$ true.

We now construct the announced finite characterizations. Let $\mathcal{O}$ be a left-atomic ontology. Choose a set Γ of $\mathcal{EL}$ concepts C with $\text{sig}(C) \subseteq \text{sig}(\mathcal{O})$ and $\text{rd}(C) \leq \text{rd}(\mathcal{O})$ such that the following two conditions are satisfied:

1. if $C_1, C_2 \in \Gamma$ are distinct, then $\not\models C_1 \equiv C_2$;
2. for every $\mathcal{EL}$ concept C with $\text{sig}(C) \subseteq \text{sig}(\mathcal{O})$ and $\text{rd}(C) \leq \text{rd}(\mathcal{O})$, there is a $C' \in \Gamma$ with $\models C \equiv C'$.

Then Γ is finite because there are only finitely many pairwise non-equivalent $\mathcal{EL}$ concepts of bounded role depth and fixed signature. We may thus set

$$P = \{(\mathcal{I}_C, a_C) \mid C \in \Gamma \text{ and } (\mathcal{I}_C, a_C) \models \mathcal{O}\} \text{ and} \\ N = \{(\mathcal{I}_C, a_C) \mid C \in \Gamma \text{ and } (\mathcal{I}_C, a_C) \not\models \mathcal{O}\}.$$

In the appendix we show that (P, N) is a finite characterization of $\mathcal{O}$. The central idea is to use Lemma 21 to show that every left-atomic $\mathcal{EL}$ ontology $\mathcal{O}'$ that fits P must be of role depth at most $\text{rd}(\mathcal{O})$ and satisfy $\mathcal{O} \models \mathcal{O}'$. If $\mathcal{O}'$ additionally fits N , Lemma 21 can be applied in the reverse direction to show that $\mathcal{O}' \models \mathcal{O}$. Thus, $\mathcal{O}$ is equivalent to $\mathcal{O}'$.

5. Conclusion

We launched the first investigation of the finite characterizability of ontologies and concept inclusions by considering the lightweight description logic $\mathcal{EL}$. Our results show that pointed examples are strictly more expressive than non-pointed examples, and can even finitely characterize the relevant class of left-atomic ontologies. We believe they are also particularly interesting from a practical perspective, as they enable smaller examples. This is because existential restrictions on the right-hand side of concept inclusions only need to be satisfied at the distinguished point.

As mentioned in the introduction, all negative results in this paper imply that, in the same setting, exact learning using membership queries alone is impossible. Conversely, for the cases in which we obtain positive results, it would be interesting to investigate whether exact learning is possible with polynomially many membership queries.

As further work, it would be interesting to investigate finite characterizability for other ontology languages such as $\mathcal{EL}_\perp$, $\mathcal{ELI}$, and (restricted) existential rules. It would also be interesting to study the complexity of deciding whether, given a set of pointed examples, there is a fitting ontology.

Acknowledgments

This work is partly supported by BMFTR (Federal Ministry of Research, Technology and Space) in DAAD project 57616814 (SECAI, School of Embedded Composite AI) as part of the program Konrad Zuse Schools of Excellence in Artificial Intelligence.

The research reported in this paper has been supported by the German Research Foundation DFG, as part of Collaborative Research Center (Sonderforschungsbereich) 1320 Project-ID 329551904 “EASE - Everyday Activity Science and Engineering”, University of Bremen (<http://www.easecrc.org/>).

Carsten Lutz was supported by DFG project LU 1417/4-1.

Declaration on Generative AI

During the preparation of this work, the author(s) used ChatGPT in order to: Grammar and spelling check, Paraphrase and reword. After using these tool(s)/service(s), the author(s) reviewed and edited the content as needed and take(s) full responsibility for the publication’s content.

References

- [1] M. M. Zloof, Query by example, Proc. of AFIPS NCC (1975).
- [2] D. M. L. Martins, Reverse engineering database queries from examples: State-of-the-art, challenges, and research opportunities, Information Systems 83 (2019) 89–100.
- [3] A. Cropper, S. Dumancic, R. Evans, S. H. Muggleton, Inductive logic programming at 30, Mach. Learn. 111 (2022) 147–172.
- [4] A. Pnueli, R. Rosner, On the synthesis of a reactive module, Proc. of POPL (1989).
- [5] J. Lehmann, P. Hitzler, Concept learning in description logics using refinement operators, Mach. Learn. 78 (2010) 203–250.
- [6] M. Funk, J. C. Jung, C. Lutz, H. Pulcini, F. Wolter, Learning description logic concepts: When can positive and negative examples be separated?, Proc. of IJCAI (2019) 1682–1688.
- [7] M. Funk, J. C. Jung, C. Lutz, Actively learning concepts and conjunctive queries under elr-ontologies, Proc. of AAAI (2021).
- [8] M. Funk, M. Grosser, C. Lutz, Fitting description logic ontologies to ABox and query examples, Proc. of KR (2025).
- [9] S. Hosemann, J. C. Jung, C. Lutz, S. Rudolph, Fitting ontologies and constraints to relational structures, Proc. of KR (2025).
- [10] M. Grosser, C. Lutz, Fitting Horn DL ontologies to ABox and query examples: A tale of simulation quantifiers and finite models, Proc. of KR (2026). To appear.
- [11] V. Gutiérrez-Basulto, J. C. Jung, L. Sabellek, Reverse engineering queries in ontology-enriched systems: The case of expressive Horn description logic ontologies, Proc. of IJCAI (2018).
- [12] J. C. Jung, C. Lutz, H. Pulcini, F. Wolter, Logical separability of labeled data examples under ontologies, Artif. Intell. 313 (2022) 103785.
- [13] B. ten Cate, V. Dalmau, Conjunctive queries: Unique characterizations and exact learnability, ACM Trans. Database Syst. 47 (2022) 14:1–14:41. doi:10.1145/3559756.
- [14] B. ten Cate, R. Koudijs, A. Ozaki, On the power and limitations of examples for description logic concepts, Proc. of IJCAI (2024) 3567–3575. doi:10.24963/ijcai.2024/395.
- [15] M. Fortin, B. Konev, V. Ryzhikov, Y. Savateev, F. Wolter, M. Zakharyashev, Unique characterisability and learnability of temporal instance queries, Proc. of KR (2022).
- [16] B. ten Cate, D. Fisman, R. Ohayon, P. Sestic, Characterizing LTL formulas by examples, arXiv preprint arXiv:2604.22097 (2026).
- [17] B. Alexe, B. ten Cate, P. G. Kolaitis, W. C. Tan, Characterizing schema mappings via data examples, ACM Trans. Database Syst. 36 (2011) 23:1–23:48.
- [18] B. ten Cate, R. Koudijs, Characterising modal formulas with examples, ACM Trans. Comput. Log. 25 (2024) 12:1–12:27.
- [19] M. Funk, J. C. Jung, C. Lutz, Frontiers and exact learning of ELI queries under DL-Lite ontologies, Proc. of IJCAI (2022).
- [20] B. Konev, C. Lutz, A. Ozaki, F. Wolter, Exact learning of lightweight description logic ontologies, J. Mach. Learn. Res. 18 (2018) 201:1–201:63.
- [21] F. Baader, I. Horrocks, C. Lutz, U. Sattler, An Introduction to Description Logic, Cambridge University Press, 2017.
- [22] J. C. Jung, C. Lutz, F. Wolter, Least general generalizations in description logic: Verification and existence, Proc. of AAAI (2020).
- [23] F. Kriegel, The distributive, graded lattice of $\mathcal{EL}$ concept descriptions and its neighborhood relation, Proc. of CLA (2018).
- [24] J. L. Ramírez Alfonsín, The Diophantine Frobenius Problem, Oxford University Press, 2005.
- [25] G. Kapetanakis, I. Rizos, An inductive proof of the Frobenius coin problem of two denominators, arXiv preprint arXiv:2308.03050 (2023).