

ABox Abduction for Inconsistent Knowledge Bases under Repair Semantics

Anselm Haak¹, Patrick Koopmann², Yasir Mahmood³, Anni-Yasmin Turhan¹

¹Knowledge Representation Group, Paderborn University, Germany

²Knowledge in Artificial Intelligence, Vrije Universiteit Amsterdam, The Netherlands

³Data Science Group, Heinz Nixdorf Institute, Paderborn University, Germany

anselm.haak@uni-paderborn.de, p.k.koopmann@vu.nl, yasir.mahmood@uni-paderborn.de,
turhan@uni-paderborn.de

Abstract

Given a knowledge base (KB) with a non-entailed fact, the ABox abduction problem asks for possible extensions of the KB that would entail this fact. This problem has many applications, ranging from diagnosis to explainability and repair. ABox abduction has been well-investigated for consistent KBs and classical semantics, but little is known for the case of inconsistent KBs, which can be caused by erroneous data. In this paper we define suitable notions of abduction in this setting and propose criteria that guide abduction towards “useful” hypotheses. To regain meaningful reasoning in the presence of inconsistencies, we use well-established repair semantics. We provide a comprehensive landscape of the complexity of ABox abduction under repair semantics, treating different variants of the abduction problem for the light-weight description logics DL-Lite and $\mathcal{EL}_\perp$.

1 Introduction

In computational logic, abductive reasoning can be defined as follows: Given two formulas K (the *background knowledge*) and Φ (the *observation*), we are looking for a *hypothesis* H s.t. $K \wedge H \models \Phi$. The original idea dates back to (Peirce 1878) as a type of non-monotonic reasoning towards finding a plausible explanation H for an unexpected observation Φ , given our knowledge K about the situation. Since then, it has been suggested for various additional use cases: to provide possible explanations for black box classifiers in *machine learning* (Ignatiev, Narydytska, and Marques-Silva 2019), to *explain missing entailments* from knowledge bases (“if H was in your knowledge base K , Φ would be entailed”) (Du, Wang, and Shen 2015; Alrabbaa et al. 2024), and to suggest possible solutions for *repairing incomplete knowledge bases* (Wei-Kleiner, Dragisic, and Lambrix 2014; Du, Wan, and Ma 2017; Haifani et al. 2022). Diagnosis, explainability and repair are well-motivated in the area of ABox reasoning with description logic (DL) ontologies (Schlobach and Cornet 2003; Baader et al. 2022): here, the ABox $\mathcal{A}$ contains a set of facts like in a database, which is used together with an ontology or TBox $\mathcal{T}$ to derive new implicit facts. *ABox abduction* is the special case of abduction in which K is such a DL knowledge base $\langle \mathcal{T}, \mathcal{A} \rangle$, Φ consists of facts, and we are looking for a set $\mathcal{H}$ of facts as hypothesis that would ensure $\langle \mathcal{T}, \mathcal{A} \cup \mathcal{H} \rangle \models \Phi$ (Klarman, Endriss, and

Schlobach 2011; Calvanese et al. 2013; Ceylan et al. 2020; Koopmann 2021). To avoid nonsensical explanations, one usually additionally requires *consistency* ($\langle \mathcal{T}, \mathcal{A} \cup \mathcal{H} \rangle \not\models \perp$), *non-triviality* ($\mathcal{H} \neq \Phi$) and sometimes puts restrictions on the *signature* of $\mathcal{H}$ (Elsenbroich, Kutz, and Sattler 2006; Koopmann et al. 2020).

Example 1. Consider the following excerpt from a medical ontology $\mathcal{T}$ on diabetes:

$$\begin{aligned} \text{High} \sqcap \text{Low} &\sqsubseteq \perp & \exists \text{glucoseLevel.High} &\sqsubseteq \text{GlycemicCrisis} \\ & & \exists \text{glucoseLevel.Low} &\sqsubseteq \text{GlycemicCrisis} \\ \exists \text{glucoseLevel.High} \sqcap \text{OverdosedInsulin} &\sqsubseteq \text{DiabeticComa} \\ \text{GlycemicCrisis} \sqcap \text{Ketoacidosis} &\sqsubseteq \text{DiabeticComa} \end{aligned}$$

The following ABox $\mathcal{A}_1$ contains medical evidence about a patient obtained through a glucose monitor:

$$\text{glucoseLevel}(\text{patient}, l) \quad \text{High}(l)$$

We observe that the patient passed out. Based on his glucose reading, a reasonable explanation is that he took too much insulin. Indeed, using ABox abduction for the observation $\text{DiabeticComa}(\text{patient})$, a possible hypothesis would be $\{\text{OverdosedInsulin}(\text{patient})\}$.

If $\langle \mathcal{T}, \mathcal{A} \rangle$ is inconsistent, every fact is entailed, and under standard semantics, both abductive and deductive reasoning cannot be used any more to make useful inferences. Unfortunately, consistency cannot always be assumed in realistic settings, as data often contains erroneous information. To overcome this issue, a range of inconsistency-tolerant semantics have been proposed (Maier, Ma, and Hitzler 2013; Bienvenu and Bourgaux 2016; Bienvenu and Maniere 2026), in particular repair-based semantics, *repair semantics* for short, which are defined based on maximal consistent subsets of the ABox called *repairs*.

Example 2. The doctor uses a finger stick sensor to get additional evidence about the glucose level of the patient, which leads to the additional ABox fact $\text{Low}(l)$. The new ABox $\mathcal{A}_2$ is inconsistent and has two repairs, namely $\mathcal{A}_1$ and $\{\text{glucoseLevel}(\text{patient}, l), \text{Low}(l)\}$. Under *Brave* semantics, which considers repairs in isolation, $\text{High}(l)$, $\text{Low}(l)$ and $\text{GlycemicCrisis}(\text{patient})$ are all entailed, while under *AR* semantics, which considers all repairs at the same time, of those only $\text{GlycemicCrisis}(\text{patient})$ is entailed. We

can adapt the definition of abduction to consider those entailment relations instead of classical entailment. Then, $\{OverdosedInsulin(patient)\}$ and $\{Ketoacidosis(patient)\}$ are both hypotheses under Brave semantics, while only $\{Ketoacidosis(patient)\}$ is a hypothesis under AR semantics. Indeed, the cautious doctor should investigate this hypothesis first.

It turns out that considering repair semantics changes a few things compared to classical semantics. The consistency requirement $\langle \mathcal{T}, \mathcal{A} \cup \mathcal{H} \rangle \not\models \perp$ becomes obsolete, and we may instead require that hypotheses should not introduce new conflicts in the knowledge base (that they are *conflict-confining*), or at least focus on those hypotheses that minimise the newly introduced conflicts. Also, allowing for the observation itself as the trivial hypothesis does not as easily trivialise the problem as under classical semantics: It might not always satisfy the conditions, and there are even cases in which other hypotheses introduce less new conflicts than the observation itself would.

While entailment under repair semantics is well-investigated, and there is work on explaining missing entailments under these semantics, to the best of our knowledge, the only work that considers abduction under repair semantics is by Du, Wang, and Shen (2015), who present a practical system for computing conflict-confining hypotheses with rewritable ontology languages under IAR semantics, another variant of repair semantics. Abduction has also been investigated for paraconsistent semantics (Bienvenu, Inoue, and Kozhemiachenko 2024), a different inconsistency-tolerant semantics, but only for propositional logic. A more complete picture for ABox abduction under repair-based semantics, considering different semantics and logics, and also analysing how to appropriately define abduction in these settings, has been missing so far.

One central motivation of abduction is to explain missing entailments, and indeed different authors have investigated this problem using a different strategy: Specifically, Bienvenu, Bourgaux, and Goasdoué (2019) and Lukasiewicz, Malizia, and Molinaro (2022) explain missing entailments under repair semantics by pointing to facts that “block” the entailment by creating conflicts in parts of the ABox that would otherwise lead to the entailment. In other words, while our explanations consist of facts that would need to be added to make the entailment true, those explanations consist of facts that need to be *removed*. However, it is easy to construct examples where an abductive solution exists, while it is not possible to make the entailment true by removing axioms. Our results thus complement these works, by providing explanations in case their approaches are incomplete.

We provide the first comprehensive analysis of abduction under repair semantics. Since complexity under these semantics easily trivialises in logics that are ExpTime-hard, we consider two prominent tractable description logics, $\mathcal{EL}_\perp$ and DL-Lite, which play a central role in many large-scale knowledge bases. We consider different variants of the problem that are suitable in the context of repair semantics, for example by requiring the aforementioned conflict-confinement as well as signature-restrictions, and also inves-

tigate different optimality criteria for hypotheses. Along the way, we show different properties of hypotheses under repair semantics that help to develop a better understanding of abduction for the analysed semantics. Some of our results also apply to the consistent setting, but have to the best of our knowledge not been published before. Our results are shown in Table 1. We recently initiated this study, with preliminary work presented at the DL workshop (Haak et al. 2025). Full proof details for our results can be found in the technical report available online (Haak et al. 2026).

2 Preliminaries

For a general introduction to description logics, we refer the reader to Baader et al. (2017). We assume familiarity with computational complexity (Sipser 1997), in particular with the complexity classes NL, P, NP, coNP, Σ_2^P and Π_2^P , as well as DP, the class of decision problems representable as the intersection of a problem in NP and a problem in coNP.

2.1 The Description Logics $\mathcal{EL}_\perp$ and DL-Lite

We use the following enumerable sets of names: N_C for *concept names*, N_R for *role names*, and N_I for *individuals*. An ABox is a finite set of *concept assertions* (also known as *flat* or *simple* assertions) of the form $A(a)$ and *role assertions* of the form $r(a, b)$ for $A \in N_C$, $r \in N_R$, $a, b \in N_I$. A TBox is a finite set of *concept inclusions* of the form $C \sqsubseteq D$ for concepts C and D and role inclusions of the form $Q \sqsubseteq S$ for roles Q and S . Concept and role inclusions are also called *axioms*. A KB is a tuple $\langle \mathcal{T}, \mathcal{A} \rangle$ for a TBox $\mathcal{T}$ and ABox $\mathcal{A}$. By *signature* of a KB $\mathcal{K}$, we mean the set of (concept, role, and individual) names appearing in $\mathcal{K}$.

We next specify the syntax for $\mathcal{EL}_\perp$ and the considered DL-Lite dialects, with semantics being defined as usual, see Baader et al. (2017). For $\mathcal{EL}_\perp$ concepts, the syntax is given by the grammar:

$$C ::= C \sqcap C \mid \exists r.C \mid A \mid \top \mid \perp.$$

An $\mathcal{EL}_\perp$ KB restricts the TBox to concept inclusions over $\mathcal{EL}_\perp$ concepts.

We consider the DL-Lite dialects DL-Lite $_{\mathcal{R}}$ and DL-Lite $_{\text{core}}$. In DL-Lite $_{\mathcal{R}}$ (underlying the OWL 2 QL profile), TBoxes may contain *concept inclusions* of the form $B \sqsubseteq C$ and *role inclusions* of the form $Q \sqsubseteq S$, where B, C, Q and S are generated by the following grammar:

$$\begin{aligned} B &::= A \mid \exists Q, & C &::= B \mid \neg B, \\ Q &::= R \mid R^-, & S &::= Q \mid \neg Q, \end{aligned}$$

where $A \in N_C$ and $R \in N_R$. DL-Lite $_{\text{core}}$ restricts DL-Lite $_{\mathcal{R}}$ by disallowing role inclusions, so that only concept inclusions of the above form are allowed. Semantics for all three DLs are defined as usual. For an ABox $\mathcal{A}$ and TBox $\mathcal{T}$, we say that $\mathcal{A}$ is $\mathcal{T}$ -consistent, if $\langle \mathcal{T}, \mathcal{A} \rangle \not\models \perp$, and $\mathcal{T}$ -inconsistent otherwise. Further, we say that $\mathcal{A}$ is a $\mathcal{T}$ -support of a concept assertion α , if $\langle \mathcal{T}, \mathcal{A} \rangle \models \alpha$.

For the rest of the paper, the general term DL-Lite refers to either DL-Lite $_{\text{core}}$ or DL-Lite $_{\mathcal{R}}$. We do so since all of our results apply to both DLs: our lower bounds apply to DL-Lite $_{\text{core}}$ and our upper bounds apply to DL-Lite $_{\mathcal{R}}$.

		Existence Problem				Verification Problem				
		general	signature	conflict-conf.	non-trivial	$\leq$ -min	$\subseteq$ -min	$\subseteq_c$ -min	$\leq_c$ -min	conflict-conf.
DL-Lite	Brave	NL ^{T13}	NL ^{T16}	NL ^{T13}	NL ^{T16}	NL ^{T11}	NL ^{T17}	NL ^{T14}	NL ^{T14}	NL ^{T14}
	AR	NL ^{T13}	coNP ^{T19}	NL ^{T13}	coNP ^{T19}	coNP ^{T11}	DP ^{T20}	NL ^{T14}	NL ^{T14}	NL ^{T14}
$\mathcal{EL}_\perp$	Brave	P ^{T22}	NP ^{T23}	Σ_2^P ^{T27}	NP ^{T24}	NP ^{T21}	DP ^{T31}	Π_2^P ^{T29}	–	DP ^{T28}
	AR	coNP ^{T22}	Σ_2^P ^{T23}	coNP ^{T27}	Σ_2^P ^{T24}	coNP ^{T21}	Π_2^P ^{T31}	coNP ^{T29}	coNP ^{T30}	coNP ^{T28}

Table 1: Complexity for existence and verification of hypotheses with different properties and repair semantics. All entries are completeness results, with superscripts referencing corresponding theorems. Verifying general hypotheses has the same complexity as with $\leq$ -minimality.

2.2 Repair Semantics

If a knowledge base is inconsistent, repair semantics can “re-store” consistent versions and admit meaningful reasoning again. We focus here on ABox repairs and define these as well as two common kinds of repair semantics next.

Let $\mathcal{K} = \langle \mathcal{T}, \mathcal{A} \rangle$ be an inconsistent knowledge base. A *repair* of $\mathcal{K}$ is a $\mathcal{T}$ -consistent subset $\mathcal{R} \subseteq \mathcal{A}$ and subset-maximal with this property, i.e., there is no $\mathcal{T}$ -consistent subset $\mathcal{R}' \subseteq \mathcal{A}$ that is a strict superset of $\mathcal{R}$. The somewhat dual notion is a *conflict* or *conflict set* $\mathcal{C}$, which is a $\mathcal{T}$ -inconsistent subset of the ABox and subset-minimal with this property. We denote by $\text{Conf}(\mathcal{K})$ the set of conflicts of $\mathcal{K}$. We recall entailment under Brave (Bienvenu and Rosati 2013) and AR semantics (Lembo et al. 2015) for concept assertions α :

- $\mathcal{K} \models_{\text{Brave}} \alpha$ iff there is a repair $\mathcal{R}$ of $\mathcal{K}$ s.t. $\langle \mathcal{T}, \mathcal{R} \rangle \models \alpha$.
- $\mathcal{K} \models_{\text{AR}} \alpha$ iff $\langle \mathcal{T}, \mathcal{R} \rangle \models \alpha$ for every repair $\mathcal{R}$ of $\mathcal{K}$.

The complexity of entailment under repair semantics is well understood (Bienvenu and Bourgaux 2016): Checking entailment of concept assertions under Brave semantics is NL-complete for DL-Lite_{core} and DL-Lite_R, and NP-complete for $\mathcal{EL}_\perp$ in combined complexity, whereas under AR semantics it is coNP-complete for all three DLs.

Besides the complexity of entailment under repair semantics, both DL-Lite dialects share the following properties for DL-Lite KBs $\mathcal{K} = \langle \mathcal{T}, \mathcal{A} \rangle$ (Bienvenu and Bourgaux 2016):

- conflicts of $\mathcal{K}$ are of size 2, and
- subset-minimal $\mathcal{T}$ -supports of concept assertions α are of size 1.

3 ABox Abduction for Inconsistent KBs

The central task of abduction is to compute hypotheses. We define these for non-entailed concept assertions under repair semantics.

Definition 3. Let $\mathcal{K} = \langle \mathcal{T}, \mathcal{A} \rangle$ be an inconsistent KB, α a concept assertion (called an *observation*) and $\mathcal{S} \in \{\text{Brave}, \text{AR}\}$ such that $\mathcal{K} \not\models_{\mathcal{S}} \alpha$. Then, the pair $\langle \mathcal{K}, \alpha \rangle$ is called an *$\mathcal{S}$ -abduction problem*. A solution for such a problem, called *$\mathcal{S}$ -hypothesis*, is an ABox $\mathcal{H}$ using only individuals occurring in $\mathcal{K}$ and α s.t. $\langle \mathcal{T}, \mathcal{A} \cup \mathcal{H} \rangle \models_{\mathcal{S}} \alpha$.

Additionally, for a set Σ (signature) containing individual, concept and role names, we define the triple $\langle \mathcal{K}, \alpha, \Sigma \rangle$ to be a *Σ -restricted $\mathcal{S}$ -abduction problem*. A solution for such a problem is an ABox $\mathcal{H}$ using only symbols from Σ that is an $\mathcal{S}$ -hypothesis for $\langle \mathcal{K}, \alpha \rangle$.

For an $\mathcal{S}$ -abduction problem $\langle \mathcal{K}, \alpha \rangle$ we require that $\mathcal{K}$ is inconsistent and $\mathcal{K} \not\models_{\mathcal{S}} \alpha$. This means we consider only the so-called *promise problem*, i.e. the problem restricted to these particular inputs. The restriction aligns with the intuition that one asks for an $\mathcal{S}$ -hypothesis if it is already known that the knowledge base is inconsistent and the observation is not $\mathcal{S}$ -entailed in $\mathcal{K}$. In contrast, if we instead assume that $\mathcal{K}$ is consistent and α is not entailed by $\mathcal{K}$ under classical semantics, we obtain classical abduction problems. In this case, we call an ABox $\mathcal{H}$ *hypothesis for α under classical semantics*, if $\langle \mathcal{T}, \mathcal{A} \cup \mathcal{H} \rangle \not\models \perp$ and $\langle \mathcal{T}, \mathcal{A} \cup \mathcal{H} \rangle \models \alpha$. Note also that we do not admit fresh individuals beyond the specified signature as it is done in Koopmann (2021). This can sometimes make the problem ExpTime-hard already in the classical case, and is left for future work.

To obtain hypotheses that are meaningful for explanation purposes, one often considers additional properties of hypotheses as well as minimality criteria that yield *preferred hypotheses*. We transfer some of this terminology already defined for abduction under classical semantics, and extend it to repair semantics by additional properties and minimality notions based on conflicts. Of those, the notion we call *conflict-confining* was already used in Du, Wan, and Ma (2017). For two sets S_1, S_2 we may abbreviate $|S_1| \leq |S_2|$ by $S_1 \leq S_2$.

Definition 4. Let $\mathcal{S} \in \{\text{Brave}, \text{AR}\}$, $\langle \mathcal{K}, \alpha \rangle$ be an $\mathcal{S}$ -abduction problem, and $\preceq \in \{\subseteq, \leq\}$. An ABox $\mathcal{H}$ is

1. *conflict-confining* for $\mathcal{K} = \langle \mathcal{T}, \mathcal{A} \rangle$, provided that $\text{Conf}(\langle \mathcal{T}, \mathcal{A} \cup \mathcal{H} \rangle) = \text{Conf}(\mathcal{K})$.
- If $\mathcal{H}$ is an $\mathcal{S}$ -hypothesis for $\langle \mathcal{K}, \alpha \rangle$, we call it
2. *non-trivial*, if $\alpha \notin \mathcal{H}$,
3. *$\preceq$ -minimal*, if there is no $\mathcal{S}$ -hypothesis $\mathcal{H}'$ for $\langle \mathcal{K}, \alpha \rangle$ such that $\mathcal{H}' \prec \mathcal{H}$, and
4. *$\preceq_c$ -minimal*, if there is no $\mathcal{S}$ -hypothesis $\mathcal{H}'$ for $\langle \mathcal{K}, \alpha \rangle$ such that $\text{Conf}(\langle \mathcal{T}, \mathcal{A} \cup \mathcal{H}' \rangle) \prec \text{Conf}(\langle \mathcal{T}, \mathcal{A} \cup \mathcal{H} \rangle)$.

While Condition 2 and 3 are standard for abduction, Condition 1 and 4 adapt the idea of a hypothesis not (resp., minimally) introducing new inconsistencies to the KB, which is already inconsistent to begin with. Conflict-confinement can be equivalently defined by requiring that $\langle \mathcal{T}, \mathcal{R} \cup \mathcal{H} \rangle \not\models \perp$ for every repair $\mathcal{R}$ of $\mathcal{K}$. Furthermore, the minimality in Condition 4 generalises this notion and allows to minimally add new conflicts. Note that hypotheses introducing new conflicts can be desirable, as erroneous facts might need to implicitly be replaced by new facts, leading to conflicts when considering both together.

We use the term *subset-minimal* for $\subseteq$ -minimal and *cardinality-minimal* for $\leq$ -minimal. Further, we also consider minimal variants of hypotheses with additional properties: For example, a $\subseteq$ -minimal conflict-confining AR-hypothesis need only be $\subseteq$ -minimal among all conflict-confining AR-hypotheses.

Conflict-confinement is a general property of an ABox and does not depend on the semantics or additional properties of abduction problems. Hence, we can already establish the complexity of checking for this property here.

Lemma 5. *Given a KB $\mathcal{K}$ and an ABox $\mathcal{H}$, checking whether $\mathcal{H}$ is conflict-confining for $\mathcal{K}$ is (1) NL-complete, if $\mathcal{K}$ is a DL-Lite KB, and (2) coNP-complete, if $\mathcal{K}$ is an $\mathcal{EL}_\perp$ KB.*

Proof sketch. NL-membership for DL-Lite follows from the fact that conflicts are of size at most 2, so we can iterate over all possible conflicts in logarithmic space and verify that they are indeed fresh conflicts introduced by $\mathcal{H}$. Hardness can be shown by reduction from directed non-reachability, constructing for a directed graph $G = (V, E)$ and $s, t \in V$ the TBox

$$\mathcal{T}_{\text{unreach}} := \{A_v \sqsubseteq A_w \mid (v, w) \in E \text{ and } t \notin \{v, w\}\} \cup \{A_t \sqsubseteq A_s\} \cup \{A_v \sqsubseteq \neg A_t \mid (v, t) \in E\}.$$

This TBox encodes edges of G by concept inclusions, but replaces all edges incident to t by a single edge from t to s and disjointness of t with all of its predecessors. This ensures that existence of an s - t -path in G is equivalent to A_t being satisfiable w.r.t. $\mathcal{T}_{\text{unreach}}$, and hence to $\{A_t(a)\}$ being conflict-confining for $\langle \mathcal{T}_{\text{unreach}}, \emptyset \rangle$.

For membership in case of $\mathcal{EL}_\perp$, we universally guess a potential conflict and verify that it is not a fresh conflict introduced by $\mathcal{H}$ in polynomial time. Hardness is obtained by a straightforward reduction from non-entailment under Brave semantics. $\square$

We investigate the following reasoning problems for a given (Σ -restricted) $\mathcal{S}$ -abduction problem.

Definition 6 (Reasoning Problems). Given an $\mathcal{S}$ -abduction problem $\langle \mathcal{K}, \alpha \rangle$ (resp. Σ -restricted $\mathcal{S}$ -abduction problem $\langle \mathcal{K}, \alpha, \Sigma \rangle$),

1. the *existence problem* asks whether $\langle \mathcal{K}, \alpha \rangle$ (resp. $\langle \mathcal{K}, \alpha, \Sigma \rangle$) has a solution, and
2. the *verification problem* asks whether a given ABox $\mathcal{H}$ is a hypothesis for $\langle \mathcal{K}, \alpha \rangle$ (resp. for $\langle \mathcal{K}, \alpha, \Sigma \rangle$).

Note that Σ -restriction and non-triviality are trivial to check for a given hypothesis, which is why we do not consider them for the complexity of verification. On the other hand, existence of minimal hypotheses is equivalent to existence of any hypothesis for all considered cases.

Under some repair semantics, already standard reasoning tasks can behave in unexpected ways. This also holds true for abduction under repair semantics, with sometimes different effects depending on the DL. For instance, there is an interesting form of non-monotonicity in the case of $\subseteq$ -minimal AR-hypotheses, which applies to both DL-Lite and $\mathcal{EL}_\perp$.

Observation 7. *Let $\langle \mathcal{K}, \alpha \rangle$ be an AR-abduction problem, where $\mathcal{K}$ is an DL-Lite or $\mathcal{EL}_\perp$ KB. Then, the set of AR-hypotheses for $\langle \mathcal{K}, \alpha \rangle$ does not need to be convex.*

Example 8. We demonstrate this by constructing a DL-Lite KB $\mathcal{K} = \langle \mathcal{T}, \mathcal{A} \rangle$ that can easily be transformed into an $\mathcal{EL}_\perp$ KB, and ABoxes $\mathcal{B}_1 \subsetneq \mathcal{B}_2 \subsetneq \mathcal{B}_3$ such that $\mathcal{B}_1$ and $\mathcal{B}_3$ are AR-hypotheses for $\langle \mathcal{K}, A(a) \rangle$, but $\mathcal{B}_2$ is not, as follows:

$$\begin{aligned} \mathcal{T} &:= \{B_1 \sqsubseteq \neg B_2, C_1 \sqsubseteq \neg C_2, B_1 \sqsubseteq A, B_3 \sqsubseteq A\}, \\ \mathcal{A} &:= \{C_1(a), C_2(a)\}, & \mathcal{B}_1 &:= \{B_1(a)\}, \\ \mathcal{B}_2 &:= \mathcal{B}_1 \cup \{B_2(a)\}, & \mathcal{B}_3 &:= \mathcal{B}_2 \cup \{B_3(a)\}. \end{aligned}$$

The TBox $\mathcal{T}$ can readily be transformed into a $\mathcal{EL}_\perp$ KB by changing the disjointness axioms to the appropriate form.

This observation implies that for $\subseteq$ -minimality, it is generally not sufficient to *locally* check all subsets that remove one assertion at a time. This leads to Π_2^P -hardness for verification of $\subseteq$ -minimal AR-hypotheses in $\mathcal{EL}_\perp$ (see Theorem 31). In contrast, we can circumvent the need for a global check in DL-Lite by building hypotheses bottom-up from singleton sets, using that in this case $\mathcal{T}$ -supports are always of size 1. This results in DP-completeness here, as shown in Lemma 18 and Theorem 20.

3.1 Effect of the Trivial Hypothesis

Because we only consider concept assertions as observations α , the ABox $\{\alpha\}$ is already a candidate hypothesis for $\langle \mathcal{K}, \alpha \rangle$ for any KB $\mathcal{K}$. We study next how such *trivial hypotheses* affect certain cases of the existence and the verification problem in our framework. Regarding Brave semantics, the following is easy to see.

Proposition 9. *Let $\langle \mathcal{K}, \alpha \rangle$ be a Brave-abduction problem. There is a Brave-hypothesis for $\langle \mathcal{K}, \alpha \rangle$ iff $\{\alpha\}$ is a Brave-hypothesis for $\langle \mathcal{K}, \alpha \rangle$.*

Proof. $\{\alpha\}$ is a Brave-hypothesis for $\langle \mathcal{K}, \alpha \rangle$ if $\{\alpha\}$ is $\mathcal{T}$ -consistent, as in this case α is contained in some repairs of $\langle \mathcal{T}, \mathcal{A} \cup \{\alpha\} \rangle$. Otherwise, if $\{\alpha\}$ is $\mathcal{T}$ -inconsistent, there is no Brave-hypothesis for $\langle \mathcal{K}, \alpha \rangle$. $\square$

We next show that for existence of general AR-hypotheses, it is sufficient to check the trivial hypothesis. Furthermore, for the trivial hypothesis, the properties of being an AR-hypothesis and being conflict-confining coincide. Note that the trivial hypothesis $\{\alpha\}$ being conflict-confining for the KB $\mathcal{K}$ means that there is *no* repair $\mathcal{R}$ of $\mathcal{K}$ s.t. $\mathcal{R} \cup \{\alpha\}$ is $\mathcal{T}$ -inconsistent. Intuitively, this can be seen as *non-entailment* of the negation of α under Brave semantics. This provides some intuition for why the complexity of checking whether an ABox is conflict-confining for a KB has the same complexity as the complement of Brave-entailment for both DLs, cf. Lemma 5.

Lemma 10. *Let $\langle \mathcal{K}, \alpha \rangle$ be an AR-abduction problem with $\mathcal{K} = \langle \mathcal{T}, \mathcal{A} \rangle$. The following are equivalent:*

1. *there is an AR-hypothesis for $\langle \mathcal{K}, \alpha \rangle$,*
2. *$\{\alpha\}$ is an AR-hypothesis for $\langle \mathcal{K}, \alpha \rangle$, and*
3. *$\{\alpha\}$ is conflict-confining for $\mathcal{K}$.*

Proof sketch. (2) $\Rightarrow$ (1) is obvious.

(3) $\Rightarrow$ (2): For any repair $\mathcal{R}$ of $\langle \mathcal{T}, \mathcal{A} \cup \{\alpha\} \rangle$, we argue that $\alpha \in \mathcal{R}$ and hence α is AR-entailed: If instead $\alpha \notin \mathcal{R}$, then $\mathcal{R}$ is a repair of $\mathcal{K}$, so $\mathcal{R} \cup \{\alpha\}$ is $\mathcal{T}$ -consistent contradicting maximality of $\mathcal{R}$.

(1) $\Rightarrow$ (3): If $\{\alpha\}$ is not conflict-confining for $\mathcal{K}$, there is a conflict of $\langle \mathcal{T}, \mathcal{A} \cup \{\alpha\} \rangle$ containing α . This yields a repair $\mathcal{R}$ of $\langle \mathcal{T}, \mathcal{A} \cup \{\alpha\} \rangle$ not containing α . Now for any ABox $\mathcal{H}$, there is a repair $\mathcal{R}'$ of $\langle \mathcal{T}, \mathcal{A} \cup \mathcal{H} \rangle$ with $\mathcal{R} \subseteq \mathcal{R}'$. But since $\mathcal{R} \cup \{\alpha\}$ is $\mathcal{T}$ -inconsistent, we have $\langle \mathcal{T}, \mathcal{R}' \rangle \not\models \alpha$, so $\mathcal{H}$ is not an AR-hypothesis for α . $\square$

The above properties simplify the corresponding existence problems, and guarantee that $\leq$ -minimal AR-hypotheses are of size 1. For AR semantics, existence of conflict-confining hypotheses is similarly affected, while for Brave semantics, the existence even becomes trivial if the concept name in α is satisfiable w.r.t. $\mathcal{T}$.

4 The Case of DL-Lite

We establish the complexity results for abduction problems over DL-Lite KBs, and establish a number of interesting properties of these problems and the corresponding hypotheses. An overview of the results is found in Table 1, whereas this section is structured by the involved proof techniques of the results. We begin by noting that verification of (general) $\mathcal{S}$ -hypotheses is essentially equivalent to $\mathcal{S}$ -entailment. By Proposition 9 and Lemma 10, this extends to $\leq$ -minimal hypotheses under both repair semantics.

Theorem 11. *Verification of general and of $\leq$ -minimal $\mathcal{S}$ -hypotheses is (1) NL-complete for $\mathcal{S} = \text{Brave}$, and (2) coNP-complete for $\mathcal{S} = \text{AR}$.*

Proof sketch. The upper bounds are obtained from those for $\mathcal{S}$ -entailment, as argued above. The lower bound for Brave semantics follows by a reduction from directed reachability, based on the following construction. Given a directed graph $G = (V, E)$ and $s, t \in V$, define

$$\mathcal{T} := \{A_v \sqsubseteq A_w \mid (v, w) \in E\}, \quad \mathcal{H}_{\text{reach}} := \{A_s(a)\}.$$

Now there is an s - t -path in G iff $\langle \mathcal{T}, \mathcal{H}_{\text{reach}} \rangle \models A_t(a)$. Adding a dummy inconsistency to obtain a KB $\mathcal{K}_{\text{reach}}$ yields the desired reduction. Note that this is a simpler variant of the construction used in the proof of Lemma 5.

The lower bound for general AR-hypotheses follows by reduction from unsatisfiability, adapting a construction from Bienvenu, Bourgaux, and Goasdoué (2019). For a CNF $\varphi(x_1, \dots, x_n) := \{c_1, \dots, c_k\}$, set $\mathcal{K}_{\text{unsat}} := \langle \mathcal{T}, \mathcal{A} \rangle$ where

$$\begin{aligned} \mathcal{T} &:= \{\exists U \sqsubseteq A, \exists P^- \sqsubseteq \neg \exists N^-\} \cup \\ &\quad \{\exists P \sqsubseteq \neg \exists U^-, \exists N \sqsubseteq \neg \exists U^-\}, \\ \mathcal{A} &:= \{P(c_j, x_i) \mid x_i \in c_j\} \cup \{N(c_j, x_i) \mid \neg x_i \in c_j\}, \text{ and} \\ \mathcal{H} &:= \{U(a, c_j) \mid j \leq k\}. \end{aligned}$$

It is known that $\langle \mathcal{T}, \mathcal{A} \cup \mathcal{H} \rangle \models_{\text{AR}} A(a)$ iff φ is unsatisfiable, yielding the desired reduction.

For $\leq$ -minimal AR-hypotheses, we use the same construction, but first transform φ into a normal form, where φ has the property that removing the last clause c from φ

always yields a satisfiable formula. This then allows us to obtain a singleton hypothesis candidate in the construction, which ensures $\leq$ -minimality. Given a formula $\varphi = \{c_1, \dots, c_k\}$ with clauses c_i , the normal form can be achieved by letting $\varphi' := \{c'_1, \dots, c'_{k+2}\}$ with

$$\begin{aligned} c'_i &:= c_i \cup \{x_{n+1}\} \text{ for } 1 \leq i \leq k, \\ c'_{k+1} &:= \neg x_{n+1} \vee x_{n+2}, \text{ and} \\ c'_{k+2} &:= \neg x_{n+2} \end{aligned}$$

We observe that φ' and φ are equisatisfiable, but removing the clause c'_{k+2} from φ' yields a satisfiable sub-formula.

Now for a formula φ in our normal form with last clause c , let $\mathcal{T}$ and $\mathcal{A}$ be as above, and define

$$\mathcal{K}' := \langle \mathcal{T}, \mathcal{A} \cup \{U(a, c') \mid c' \in \varphi \setminus \{c\}\} \rangle$$

Here, $\mathcal{H}' := \{U(a, c)\}$ is an AR-hypothesis for $\langle \mathcal{K}', A(a) \rangle$ iff φ is unsatisfiable, since $\varphi \setminus \{c\}$ is never unsatisfiable. As $|\mathcal{H}'| = 1$, $\leq$ -minimality is ensured. $\square$

Turning towards the existence problem, Proposition 9 and Lemma 10 show that for general Brave-hypotheses as well as for general and for conflict-confining AR-hypotheses, it is sufficient to check the trivial hypothesis. We next show that for DL-Lite, the same applies to conflict-confining Brave-hypotheses by showing that here, Brave-hypotheses have a very simple structure.

Lemma 12. *Let $\langle \mathcal{K}, \alpha \rangle$ be a Brave-abduction problem, where $\mathcal{K} = \langle \mathcal{T}, \mathcal{A} \rangle$ is a DL-Lite KB, and assume there is a Brave-hypothesis $\mathcal{H}$ for $\langle \mathcal{K}, \alpha \rangle$. Then there is an assertion $\beta \in \mathcal{H}$ s.t. $\{\beta\}$ is a Brave-hypothesis for $\langle \mathcal{K}, \alpha \rangle$. Furthermore, $\text{Conf}(\langle \mathcal{T}, \mathcal{A} \cup \{\beta\} \rangle) \subseteq \text{Conf}(\langle \mathcal{T}, \mathcal{A} \cup \mathcal{H} \rangle)$ and for any conflict $\{\alpha, \gamma\} \in \text{Conf}(\langle \mathcal{T}, \mathcal{A} \cup \{\alpha\} \rangle)$, we have $\{\beta, \gamma\} \in \text{Conf}(\langle \mathcal{T}, \mathcal{A} \cup \{\beta\} \rangle)$.*

In particular, the last part of the lemma shows that existence of any conflict-confining Brave-hypothesis is equivalent to $\{\alpha\}$ being one. Now, using the upper bound for checking that an ABox is conflict-confining in Lemma 5, we obtain following complexity results. Surprisingly, the complexity under AR semantics drops below even that of AR entailment. This can be explained by the relationship to brave non-entailment noted before Lemma 10.

Theorem 13. *The existence problems for general and for conflict-confining $\mathcal{S}$ -hypotheses are NL-complete for $\mathcal{S} \in \{\text{Brave}, \text{AR}\}$.*

Proof sketch. For general Brave-hypotheses, membership readily follows from Proposition 9. The remaining cases are equivalent to checking that $\{\alpha\}$ is conflict-confining by Lemmata 10 and 12, and the complexity follows from Lemma 5. Hardness for Brave-hypotheses can be shown by a similar reduction as the one used in Lemma 5, adding a dummy inconsistency. $\square$

Building on these techniques, we next show that verification of conflict-confining and conflict-minimal hypotheses enjoys the same low complexity.

Theorem 14. *Verification of conflict-confining and of $\preceq_c$ -minimal $\mathcal{S}$ -hypotheses is NL-complete for $\preceq \in \{\subseteq, \leq\}$ and $\mathcal{S} \in \{\text{Brave}, \text{AR}\}$.*

Proof sketch. Hardness for all cases follows by reduction from directed reachability as in the proof of Theorem 11, noting that $\mathcal{H}_{\text{reach}}$ is conflict-confining for $\mathcal{K}_{\text{reach}}$. To verify that $\mathcal{H}$ is a conflict-confining Brave-hypothesis, verify separately that it is a Brave-hypothesis and that it is conflict-confining, each possible in NL. For the latter, use that conflicts are of size at most 2, similar to Theorem 13. For AR, first note that here, a hypothesis is $\preceq_c$ -minimal iff it is conflict-confining by Lemma 10, so we can focus on the conflict-confining case. We show that conflict-confining hypotheses have a simpler structure than general ones: together with the TBox, such a hypothesis already entails α classically, leading to an NL-algorithm. For the case of $\preceq_c$ -minimal Brave-hypotheses, we first look for a singleton hypothesis $\{\beta\} \subseteq \mathcal{H}$, which must exist in any hypothesis by Lemma 12. This simplifies the minimality-checks, yielding NL-algorithms for both cases by carefully comparing conflicts for $\{\alpha\}$ and $\{\beta\}$ and using Lemma 12. $\square$

The remaining cases under Brave semantics behave similarly, with membership again following from Lemma 12. In case of existence of non-trivial hypotheses for some observation $A(a)$, one may be tempted to only consider the direct subsumees of A at least for Brave semantics. Notably, this is not sufficient, as demonstrated by the following example, which is also incorporated in the construction for hardness.

Example 15. Define

$$\mathcal{T} := \{B \sqsubseteq \exists r, \exists r \sqsubseteq A, A \sqsubseteq \neg \exists r^-, C \sqsubseteq \neg C\},$$

and let $\mathcal{K} := \langle \mathcal{T}, \{C(a)\} \rangle$. Then $\langle \mathcal{K}, A(a) \rangle$ is a Brave-abduction problem, as $\mathcal{K}$ is obviously inconsistent and $\mathcal{K} \not\models_{\text{Brave}} A(a)$. The ABox $\mathcal{B} := \{r(a, a)\}$, based on the direct subsumee $\exists r$ of A , is not a Brave-hypothesis for $\langle \mathcal{K}, A(a) \rangle$, as $\langle \mathcal{T}, \mathcal{B} \rangle \models \perp$. Still, there is such a hypothesis, namely $\mathcal{H} := \{B(a)\}$.

Intuitively, not allowing fresh individuals in a hypothesis has the effect that we have to find a subsumee satisfiable with the limited number of individuals present in the KB. Note that the last TBox axiom and the assertion $C(a)$ in the ABox are only necessary to ensure inconsistency, and do not interfere with the rest of the construction.

Theorem 16. *The existence problems for Σ -restricted and for non-trivial Brave-hypotheses are NL-complete.*

Theorem 17. *Verification of $\subseteq$ -minimal Brave-hypotheses is NL-complete.*

The remaining cases of existence of AR-hypotheses would seem to require guessing a hypothesis and verifying using AR-entailment, which would result in a Σ_2^P -algorithm. At first, this seems further evidenced by the non-convexity of AR-hypotheses, seen in Observation 7: For example, for $\subseteq$ -minimality, it is not sufficient to check all direct subsets only removing one element. Surprisingly, we can still get around this obstacle, due to the following property.

Lemma 18. *Let $\langle \mathcal{K}, \alpha \rangle$ be an AR-abduction problem, where $\mathcal{K} = \langle \mathcal{T}, \mathcal{A} \rangle$ is a DL-Lite KB, and let $\mathcal{B}$ be a set of assertions. There is an AR-hypothesis $\mathcal{H} \subseteq \mathcal{B}$ for $\langle \mathcal{K}, \alpha \rangle$ iff for each repair $\mathcal{R}$ of $\mathcal{K}$, there is a $\mathcal{T}$ -support $\{\beta\} \subseteq \mathcal{A} \cup \mathcal{B}$ of α s.t. $\mathcal{R} \cup \{\beta\}$ is $\mathcal{T}$ -consistent.*

Proof sketch. ($\Leftarrow$) Let $\mathcal{R}_1, \dots, \mathcal{R}_m$ be the repairs of $\mathcal{K}$ and $\beta_1, \dots, \beta_m$ the $\mathcal{T}$ -supports of α , where $\mathcal{R}_i \cup \{\beta_i\}$ is $\mathcal{T}$ -consistent. Each repair $\mathcal{R}$ of $\langle \mathcal{T}, \mathcal{A} \cup \{\beta_1, \dots, \beta_m\} \rangle$ must contain at least one of the β_i , and hence a $\mathcal{T}$ -support of α : Otherwise, $\mathcal{R} \subseteq \mathcal{A}$ and $\mathcal{R} \cup \{\beta_j\}$ must be $\mathcal{T}$ -consistent for some j , leading to contradiction.

($\Rightarrow$) Any repair $\mathcal{R}$ of $\mathcal{K}$ is contained in some repair $\mathcal{R}'$ of $\langle \mathcal{T}, \mathcal{A} \cup \mathcal{H} \rangle$. As $\mathcal{H}$ is an AR-hypothesis for α , there is some $\mathcal{T}$ -support $\{\beta\} \subseteq \mathcal{R}'$ of α . As $\mathcal{R} \cup \{\beta\} \subseteq \mathcal{R}'$ and $\mathcal{R}'$ is a repair, $\mathcal{R} \cup \{\beta\}$ is $\mathcal{T}$ -consistent. $\square$

We use this to find hypotheses $\mathcal{H} \subseteq \mathcal{B}$ in coNP. This resolves the remaining cases, starting with the remaining existence problems.

Theorem 19. *The existence problems for Σ -restricted and for non-trivial AR-hypotheses are coNP-complete.*

Proof sketch. The following is a coNP-algorithm for existence of Σ -restricted AR-hypotheses by Lemma 18: Given a Σ -restricted AR-abduction problem $\langle \mathcal{K}, \alpha, \Sigma \rangle$, let $\mathcal{B}$ be the set of assertions over Σ . Guess a repair of $\mathcal{K}$ and check for a $\mathcal{T}$ -support $\{\beta\} \subseteq \mathcal{A} \cup \mathcal{B}$ s.t. $\mathcal{R} \cup \{\beta\}$ is $\mathcal{T}$ -consistent. Both checks are in NL $\subseteq$ P. The non-trivial case can be handled with the same algorithm, only changing $\mathcal{B}$ to be the set of assertions over the signature of $\langle \mathcal{K}, \alpha \rangle$, except for α itself.

Hardness for existence of Σ -restricted AR-hypotheses is shown by adapting the construction used in Theorem 11: we choose the signature $\Sigma := \{U, a, c_j \mid j \leq k\}$ and prevent the only remaining unintended hypothesis $\{U(a, a)\}$ by adding the axiom $\exists U \sqsubseteq \neg \exists U^-$.

To show hardness for non-trivial AR-hypotheses, we instead reduce from checking whether a given $\forall$ -QBF $\forall X \varphi(X)$ is true, where $\varphi = \{C_1, \dots, C_k\}$ is in DNF. We encode φ into the KB $\mathcal{K} = \langle \mathcal{T}, \mathcal{A} \rangle$ as follows:

$$\begin{aligned} \mathcal{T} := & \{C_j \sqsubseteq A_\varphi \mid 1 \leq j \leq m\} \cup \\ & \{T_{x_i} \sqsubseteq \neg C_j \mid \neg x_i \in C_j\} \cup \\ & \{F_{x_i} \sqsubseteq \neg C_j \mid x_i \in C_j\} \cup \\ & \{T_{x_i} \sqsubseteq \neg F_{x_i} \mid 1 \leq i \leq n\} \text{ and} \\ \mathcal{A} := & \{T_{x_i}(a), F_{x_i}(a) \mid 1 \leq i \leq n\}. \end{aligned}$$

Here, $\mathcal{T}$ encodes that φ is true, if at least one clause C_j is true, and a clause C_j is falsified when at least one of its literals is falsified. It remains to show that $\forall X \varphi(X)$ is true iff there is a non-trivial AR-hypothesis for $\langle \mathcal{K}, A_\varphi(a) \rangle$. For this, note that the repairs of $\mathcal{K}$ correspond to assignments over X , and that the only non-trivial assertions that can help entailment of $A(a)$ are assertions of the form $C_j(a)$. $\square$

Even though verification of $\subseteq$ -minimal AR-hypotheses has higher complexity than the previous cases, the upper bound again relies on Lemma 18.

Theorem 20. *Verification of $\subseteq$ -minimal AR-hypotheses is DP-complete.*

Proof sketch. To verify a $\subseteq$ -minimal AR-hypothesis $\mathcal{H} = \{\beta_1, \dots, \beta_n\}$, one needs to check that it is an AR-hypothesis in coNP, and that it is $\subseteq$ -minimal. For the latter, define the direct subsets $\mathcal{B}_i := \mathcal{H} \setminus \{\beta_i\}$, noting that $\mathcal{H}$ is not $\subseteq$ -minimal iff there is any AR-hypothesis $\mathcal{H}' \subseteq \mathcal{B}_i$ for some i . Hence, we can check non-minimality in coNP using n subsequent runs of the algorithm from the membership proof of Theorem 19, replacing $\mathcal{B}$ by the different $\mathcal{B}_i$. This yields an NP-algorithm for $\subseteq$ -minimality, and DP-membership in total. For hardness, we reuse the construction in the proof of Theorem 11. Specifically, we reduce from the known DP-hard problem of checking whether a set of clauses $\psi \subseteq \varphi$ is a *minimal unsatisfiable subset* (MUS) of φ (Liberatore 2005). $\square$

5 The Case of $\mathcal{EL}_\perp$

For complexity in $\mathcal{EL}_\perp$, we begin with cases that behave analogous to the corresponding cases for DL-Lite. Verification of general hypotheses is again essentially equivalent to entailment for both semantics. Moreover, Proposition 9 and Lemma 10 allow us to extend results to the $\leq$ -minimal case and to handle existence of general hypotheses.

Theorem 21. *Verification of general $\mathcal{S}$ -hypotheses and $\leq$ -minimal $\mathcal{S}$ -hypotheses is (1) NP-complete for $\mathcal{S} = \text{Brave}$, and (2) coNP-complete for $\mathcal{S} = \text{AR}$.*

Theorem 22. *The existence problem for general $\mathcal{S}$ -hypotheses is (1) P-complete for $\mathcal{S} = \text{Brave}$, and (2) coNP-complete for $\mathcal{S} = \text{AR}$.*

We next address each property of hypotheses in turn. An appropriate Σ -restriction can disallow the trivial hypothesis, so that for deciding existence in this setting, reasoning for only the trivial hypothesis is not sufficient. This interestingly leads to existence of Brave-hypothesis still enjoying the same complexity as Brave-entailment, while the complexity for AR semantics jumps to the second level of the polynomial hierarchy.

Theorem 23. *The existence problem for Σ -restricted $\mathcal{S}$ -hypotheses is (1) NP-complete for $\mathcal{S} = \text{Brave}$, and (2) Σ_2^P -complete for $\mathcal{S} = \text{AR}$.*

Proof sketch. Membership for (1): Guess an ABox $\mathcal{H}$ over signature Σ as well as a candidate repair $\mathcal{R} \subseteq \mathcal{A} \cup \mathcal{H}$, and verify that $\mathcal{R}$ is $\mathcal{T}$ -consistent and $\langle \mathcal{T}, \mathcal{R} \rangle \models \alpha$ in polynomial time. Membership for (2): Guess an ABox $\mathcal{H}$ over Σ and check that $\langle \mathcal{T}, \mathcal{A} \cup \mathcal{H} \rangle \models_{\text{AR}} \alpha$. The latter can be handled by an NP-oracle, yielding Σ_2^P -membership.

Hardness for (1) follows by reduction from SAT. Let $\varphi(X) := \{c_1, \dots, c_n\}$ be a CNF with clauses c_i . We model the satisfaction of φ by the TBox

$$\mathcal{T} := \{A_x \sqcap A_{\bar{x}} \sqsubseteq \perp \mid x \in X\} \cup \left\{ A_\ell \sqsubseteq A_c \mid \ell \in c, c \in \varphi \right\} \cup \left\{ \bigcap_{c \in \varphi} A_c \sqsubseteq A_\varphi \right\}.$$

Let $\alpha := A_\varphi(m)$ and $\Sigma := \{A_x, A_{\bar{x}} \mid x \in X\} \cup \{m\}$. The reduction actually works with an empty ABox and hence

also applies to the existence of a hypothesis over Σ in a consistent KB $\mathcal{K}$. Any hypothesis $\mathcal{H}$ over Σ corresponds to an assignment $\theta_{\mathcal{H}}$ over X , and entailment $A_\varphi(m)$ in $\langle \mathcal{T}, \mathcal{A} \cup \mathcal{H} \rangle$ is equivalent to $\theta_{\mathcal{H}} \models \varphi$.

Hardness for (2) follows by reduction from $\exists\forall$ -QBF, utilising a similar construction as in (1). Given an $\exists\forall$ -QBF $\Phi = \exists Y \forall Z. \varphi$ with φ in DNF, the main differences are that we encode satisfaction of a DNF instead of a CNF, and that the assertions encoding an assignment over Z are included in the KB, while those encoding an assignment over Y have to be part of the hypothesis. $\square$

Existence of non-trivial hypotheses turns out to have the same complexity as existence of Σ -restricted hypotheses, but for Brave semantics requires us to incorporate the idea from Example 15 into the reduction for NP-hardness.

Theorem 24. *The existence problem for non-trivial $\mathcal{S}$ -hypotheses is (1) NP-complete for $\mathcal{S} = \text{Brave}$, and (2) Σ_2^P -complete for $\mathcal{S} = \text{AR}$.*

Proof sketch. Membership in both cases is shown similar as in the proof of Theorem 23, so we focus on hardness.

For (1), we build on the construction from the proof of Theorem 23. The main change is to use the idea from Example 15 to replace the explicit signature-restriction, enforcing use of the assertions of the form $A_\ell(m)$ for a literal ℓ in any hypothesis. This can be achieved by replacing each $A_\ell \sqsubseteq A_c$ by $A_\ell \sqsubseteq \exists r_c. B$ and using the axiom $\bigcap_{c \in \varphi} \exists r_c. B \sqsubseteq A_\varphi$ to entail A_φ instead. Moreover, we add $A_\varphi \sqcap B \sqsubseteq \perp$ to trigger a conflict when trying to entail $A_\varphi(m)$ using assertions $r_c(m, m)$ and $B(m)$ instead of the intended assertions of the form $A_\ell(m)$. Note that this makes use of the fact that fresh individuals are not allowed, and that non-triviality rules out using the assertion $A_\varphi(m)$ in the hypothesis.

As before, the reduction works with an empty ABox and hence also applies to consistent KBs. Our TBox takes the following final form:

$$\mathcal{T} := \{A_x \sqcap A_{\bar{x}} \sqsubseteq \perp \mid x \in X\} \cup \left\{ A_\ell \sqsubseteq \exists r_c. B \mid \ell \in c, c \in \varphi \right\} \cup \left\{ \bigcap_{c \in \varphi} \exists r_c. B \sqsubseteq A_\varphi \right\}.$$

For (2): we reduce directly from an instance of $\exists\forall$ -QBF. Let $\Psi := \exists Y \forall Z. \psi$ be a formula with $\psi := \{t_1, \dots, t_n\}$ a DNF where each t_i is a term over $X = Y \cup Z$. We construct the KB $\mathcal{K} := \langle \mathcal{T}, \mathcal{A} \rangle$, where

$$\mathcal{T} := \{A_x \sqcap A_{\bar{x}} \sqsubseteq \perp \mid x \in X\} \cup \left\{ C \sqcap \bigcap_{\ell \in t} A_\ell \sqsubseteq A_\psi \mid t \in \psi \right\}, \text{ and} \\ \mathcal{A} := \{A_z(m), A_{\bar{z}}(m) \mid z \in Z\}.$$

Moreover, we let $\alpha := A_\psi(m)$ be our observation. Intuitively, although $\{\alpha\}$ is a trivial AR-hypotheses, any *non-trivial* AR-hypothesis corresponds to a satisfying assignment over Y for Ψ . Thus, Ψ is true iff α admits a non-trivial AR-hypothesis in $\mathcal{K}$. $\square$

Existence of conflict-confining AR-hypotheses behaves very similar as for DL-Lite, with our proofs mostly relying on Lemma 10. In sharp contrast, the setting behaves very different under Brave semantics. While this seems counter-intuitive, the following observation applies to $\mathcal{EL}_\perp$, as demonstrated in the example below.

Observation 25. *There is a Brave-abduction problem $\langle \mathcal{K}, \alpha \rangle$ s.t. there is a conflict-confining Brave-hypothesis for $\langle \mathcal{K}, \alpha \rangle$, but $\{\alpha\}$ is not such a hypothesis.*

Example 26. Let $\mathcal{K} = \langle \mathcal{T}, \mathcal{A} \rangle$, where

$$\begin{aligned} \mathcal{T} &= \{ A \sqcap B \sqsubseteq \perp, B \sqcap C \sqsubseteq \perp, C \sqcap D \sqsubseteq A \}, \text{ and} \\ \mathcal{A} &= \{ B(a), C(a) \}, \end{aligned}$$

and let $\alpha = A(a)$. It is easy to see that $\{\alpha\}$ is a Brave-hypothesis for $\langle \mathcal{K}, \alpha \rangle$, but results in a new conflict $\{A(a), B(a)\}$ of $\langle \mathcal{T}, \mathcal{A} \cup \{\alpha\} \rangle$. Hence, $\{\alpha\}$ is not conflict-confining in $\mathcal{K}$. However, $\mathcal{H} := \{D(a)\}$ entails $A(a)$ in the repair $\{C(a), D(a)\}$ and is consistent with all repairs of $\mathcal{K}$. Intuitively, the only repair of $\mathcal{A} \cup \mathcal{H}$ where α is entailed, is the one that already got rid of the conflict with α .

This contrasts with the situation for DL-Lite, where a Brave-abduction problem with this property does not exist due to Lemma 12. Surprisingly, the result is that adding the conflict-confining restriction decreases the complexity under AR semantics, while it increases it for Brave semantics. Note that this is the opposite behaviour to what was seen for Σ -restriction and non-triviality in Theorems 23 and 24, where complexity increased under AR semantics and stayed the same for Brave semantics.

Theorem 27. *The existence problem for conflict-confining S -hypotheses is (1) Σ_2^P -complete for $S = \text{Brave}$, and (2) coNP-complete for $S = \text{AR}$.*

Proof sketch. For membership in (1): although one can guess a hypothesis and a witnessing candidate repair simultaneously, the conflict-confinement of this hypothesis requires oracle calls. In contrast, the membership for (2) holds due to the observation that if there exists some conflict-confining AR-hypothesis, the observation itself must be conflict-confining and hence an AR-hypothesis. Hence, one only requires to check this later condition for the observation, which can be done in coNP due to Lemma 5.

For hardness in (1): we reduce from $\exists\forall$ -QBF reusing and adapting ideas from the reduction for AR semantics in the proof of Theorem 23. Intuitively, the outer quantifier simulates the existence of a hypothesis, whereas the inner one now checks the conflict-confinement. The required changes include (1) shifting from AR to Brave entailment, (2) causing new conflicts when a hypothesis contains an assertion for concepts outside Σ , and (3) encoding the situation when a partial assignment over Z already satisfies the formula. Here (3) is essential to encode that new conflicts are not subsumed by existing ones. $\square$

Verification of conflict-confining and $\subseteq_c$ -minimal hypotheses for can be treated using similar ideas. Again, the complexity under Brave semantics increases when requiring the hypotheses to be conflict-confining (or $\subseteq_c$ -minimal),

while it stays the same as for verification of general hypotheses in case of AR, due to the effect of the trivial hypothesis.

Theorem 28. *Verification of conflict-confining S -hypotheses is (1) DP-complete for $S = \text{Brave}$, and (2) coNP-complete for $S = \text{AR}$.*

Theorem 29. *Verification of $\subseteq_c$ -minimal S -hypotheses is (1) Π_2^P -complete for $S = \text{Brave}$, and (2) coNP-complete for $S = \text{AR}$.*

Proof sketch. We sketch the proof for (1). For membership: Given an ABox $\mathcal{H}$, we check whether $\mathcal{H}$ is a Brave-hypothesis and guess a counter-witness $\mathcal{H}'$ to $\mathcal{H}$ being $\subseteq_c$ -minimal. Verifying that either $\mathcal{H}$ is not a Brave-hypothesis, or $\mathcal{H}'$ is a Brave-hypothesis and causes fewer conflicts than $\mathcal{H}$ is done using an NP-oracle. This yields membership in Π_2^P .

For hardness: we reuse the reduction from the proof of Theorem 27. To this aim, we construct the KB $\mathcal{K}'$ from $\mathcal{K}$ in the previous construction by adding to the TBox the axioms

$$C_d \sqcap X \sqsubseteq C \quad \text{and} \quad X \sqcap Y \sqsubseteq \perp$$

and to the ABox the assertion $Y(m)$. Then, let $\alpha := C(m)$ as before and take $\mathcal{H} := \{X(m)\}$. Intuitively, $\mathcal{H}$ uses these new axioms to entail α while inducing exactly one new conflict in $\mathcal{K}'$, namely $\{X(m), Y(m)\}$. We then argue that $\mathcal{H}$ is not a $\subseteq_c$ -minimal hypotheses for $\langle \mathcal{K}, C(m) \rangle$ iff $\langle \mathcal{K}, C(m) \rangle$ admits a conflict-confining Brave-hypothesis. For this, observe that the latter is also a conflict-confining Brave-hypothesis for $\langle \mathcal{K}', C(m) \rangle$, and therefore a hypothesis introducing less conflicts than $\mathcal{H}$. $\square$

The case of $\leq_c$ -minimal Brave-hypotheses is more challenging: Here, when adapting the algorithm for the case of $\subseteq_c$ -minimality, we would have to compare the *number* of conflicts introduced by the given candidate hypothesis $\mathcal{H}$ to the *number* of conflicts introduced by a potential counter-witness. As it is not hard to see that the number of conflicts can be exponential, this would naively require an oracle for the counting class $\#P$. It seems likely that this case requires quite different techniques, and we for now only observe that the problem is certainly in PSPACE by a straightforward algorithm that counts the number of conflicts, and Π_2^P -hard by the same proof as for the case of $\subseteq_c$ -minimality above. In contrast, verification of $\leq_c$ -minimal AR-hypotheses can again be handled using similar techniques as above, due to the effect of the trivial hypothesis.

Theorem 30. *Verification of $\leq_c$ -minimal AR-hypotheses is coNP-complete.*

Finally, $\subseteq_c$ -minimality increases the complexity of verification for both semantics compared to the $\leq_c$ -minimality. In case of AR semantics, the non-convexity of the set of AR-hypotheses comes into play. In contrast to the case of DL-Lite, where Lemma 18 allowed us to circumvent this, we can here combine the idea of Example 8 with an encoding of a $\forall\exists$ -QBF to show Π_2^P -hardness using conjunction.

Theorem 31. *Verification of $\subseteq_c$ -minimal S -hypotheses is (1) DP-complete for $S = \text{Brave}$, and (2) Π_2^P -complete for $S = \text{AR}$.*

Proof sketch. Membership for both semantics is relatively straightforward. For hardness under Brave semantics, we use a reduction from the combination of an entailment (NP-hard) and a non-entailment question (coNP-hard) under Brave semantics, which is DP-hard.

Hardness for $\mathcal{S} = \text{AR}$ is more involved. Here, we reduce a $\forall\exists$ -QBF $\Phi = \exists Y \forall Z \varphi(Y, Z)$ to an AR-abduction problem $\langle \mathcal{K}, A(a) \rangle$ and ABox $\mathcal{H}$ s.t. $\mathcal{H}$ is always an AR-hypothesis of $\langle \mathcal{K}, A(a) \rangle$, and there is some AR-hypothesis $\mathcal{H}' \subsetneq \mathcal{H}$ of $\langle \mathcal{K}, A(a) \rangle$ iff Φ is true. Intuitively, subsets $\mathcal{H}' \subsetneq \mathcal{H}$ encode assignments over Y .

The construction builds on two main ideas: First, two disjoint concepts B_1 and B_2 split the set of repairs into two parts. Each repair $\mathcal{R}$ in the first part encodes an assignment over Z , and entailment of $A(a)$ in $\langle \mathcal{T}, \mathcal{R} \rangle$ is equivalent to φ being true under assignment encoded by $\mathcal{H}'$ and $\mathcal{R}$. The second part contains a single repair ensuring that $\mathcal{H}'$ encodes a full assignment over Y . Second, building on the idea from Example 8, we use an additional axiom to ensure that $\mathcal{H}$ is always an AR-hypothesis for $\langle \mathcal{K}, A(a) \rangle$ without interfering with the rest of the construction. $\square$

6 Conclusion and Outlook

We extend abduction from the classical setting to KBs that admit inconsistencies. This is achieved by designing properties that govern abductive hypotheses under commonly studied repair semantics. We present a comprehensive, detailed, and almost complete complexity landscape under all considered criteria, with only the complexity of verification for $\leq_c$ -minimal Brave-hypotheses in $\mathcal{EL}_\perp$ remaining open. In summary, abduction for DL-Lite behaves largely similarly to entailment under the corresponding semantics, with certain cases where it becomes even easier (e.g., the existence of conflict-confining AR-hypotheses). Nevertheless, membership in many cases for DL-Lite is derived using new insights for considered properties. For $\mathcal{EL}_\perp$, the situation differs, as there are cases where the complexity exceeds the one of the corresponding entailment problem under both semantics. Our complexity classification highlights how different properties of *preferred* hypotheses affect the considered (Brave and AR) semantics.

Additional findings are certain effects observed for abduction under the considered semantics. We mention two notable cases here. First, *non-convexity* (Observation 7) applies to both DLs but leads to higher complexity only for $\mathcal{EL}_\perp$, as a workaround exists for DL-Lite (see Lemma 18). Second, an observation may admit a conflict-confining Brave-hypothesis even if the observation itself is not conflict-confining for $\mathcal{EL}_\perp$ KBs (Observation 25). Regarding $\mathcal{EL}_\perp$, it is also worth noting how the complexity landscape compares to the setting of consistent KBs. While brave semantics often behaves as in the consistent case, conflict-confinement increases complexity and has no meaningful counterpart in the consistent setting. Moreover, verification in nearly all cases incurs higher complexity, in contrast to the classical setting, where the problem reduces to classical entailment and can be decided efficiently. This can be partially explained again by the non-convex nature of the space of admissible hypotheses.

We see several directions to further pursue this work. Beyond completing the picture by determining the precise complexity for the only remaining open case, an obvious next step is to study *combinations* of properties and seeing their effects on the complexity. We can already make interesting observations that further motivate this investigation. First, considering signature-restriction or non-triviality together with conflict-confinement, the complexity for both semantics jumps to the second level of the polynomial hierarchy.

Corollary 32. *For $\mathcal{EL}_\perp$, the existence problem for signature-restricted (or non-trivial) conflict-confining $\mathcal{S}$ -hypotheses is Σ_2^P -complete for $\mathcal{S} \in \{\text{Brave}, \text{AR}\}$.*

Interestingly, when considering Σ -restriction, one might expect that non-triviality can be obtained “for free” by omitting from the signature some of the symbols of the observation, disallowing the trivial hypothesis. However, a non-trivial hypothesis may reuse the symbols of the observation in a different manner, e.g., to satisfy an existential restriction involving the concept from the observation. This effect is illustrated in the following example, and motivates to also study the combination of non-triviality with Σ -restriction.

Example 33. Let $\mathcal{K} := \langle \mathcal{T}, \mathcal{A} \rangle$ be the KB with $\mathcal{T} := \{A \sqcap B \sqsubseteq C, D \sqcap \exists r.C \sqsubseteq A\}$ and $\mathcal{A} := \{B(m), r(m, n)\}$. Moreover, let $\alpha := C(m)$ and $\Sigma := \{C, D, m, n\}$. Here, the trivial hypothesis $\{C(m)\}$ is in fact also Σ -restricted, while the hypothesis $\mathcal{H}_1 := \{A(m)\}$ is non-trivial, but outside our signature. Finally, $\mathcal{H} := \{C(n), D(m)\}$ is non-trivial and Σ -restricted. It uses C and m , that is, both the concept name and the individual from the observation.

It also seems interesting to study the so-called *data complexity* of abductive reasoning. Here, one looks at the complexity purely in terms of the ABox while keeping the TBox fixed, motivated by the situation in practice, where the amount of data often far exceeds the size of the TBox. While several of our results for DL-Lite already extend to this setting, as the construction for hardness uses a TBox of constant size (e.g. coNP-hardness in Theorem 11), a systematic complexity analysis remains open for further exploration.

Allowing fresh individuals in hypotheses also seems interesting, as it may impact the complexity in certain cases. Indeed, even some of our hardness proofs rely on the fact that we do not admit fresh individuals (e.g., Theorems 22 and 24). With fresh individuals in the hypotheses, an intriguing case arises for $\mathcal{EL}_\perp$: (1) existence of non-trivial Brave-hypotheses becomes easy, since one now only needs to consider direct subsumees of the concept appearing in the observation, and (2) existence of conflict-confining hypothesis might get harder as there is no obvious polynomial bound for the size of hypotheses in this setting. Indeed, it is known that admitting fresh individuals in the signature-restricted setting may lead to exponentially large hypotheses (Koopmann 2021).

Another prominent future direction is to consider expressive observations in an abduction problem. This includes extending our analysis from *flat* to *complex* concepts in an observation, as well as to Boolean (conjunctive) queries. The hardness results from our work already transfer, whereas we

anticipate increased complexity in certain cases. Additional future work includes exploring alternative definitions of *preferred* hypotheses, such as semantically minimal ones (Du, Wang, and Shen 2015). In classical abduction, a hypothesis $\mathcal{H}$ is *semantically minimal*, if there exists no hypothesis $\mathcal{H}'$ such that $\langle \mathcal{T}, \mathcal{A} \cup \mathcal{H} \rangle \models \mathcal{H}'$, but $\langle \mathcal{T}, \mathcal{A} \cup \mathcal{H}' \rangle \not\models \mathcal{H}$. We argue that while such a minimality criterion is natural for AR semantics, its meaning is unclear for Brave-hypotheses. Further exploration of this minimality criterion is therefore left for future work.

Recently, *contrastive explanations* (Koopmann et al. 2026) have been introduced to explain why an individual (the *fact*) is entailed to belong to a concept, while another (the *foil*) is not, by identifying ABox assertions that would entail the foil. Adding such assertions may introduce conflicts and the current approach considers entailment in *some* repair that avoids all such conflicts. As future work, we propose to extend this notion using repair-based abduction, enabling contrastive explanations that explicitly account for and resolve such conflicts.

AI Declaration

The authors used generative AI to help with finding a fitting example for the introduction. Otherwise no generative AI tools were used by the authors in the production of this paper.

Acknowledgments

We thank all anonymous reviewers for their valuable feedback that helped us improve the exposition of our paper. The third author appreciates funding by the Deutsche Forschungsgemeinschaft (DFG, German Research Foundation), grant TRR 318/1 2021 – 438445824.

References

- Alrabbaa, C.; Borgwardt, S.; Friese, T.; Hirsch, A.; Knieriem, N.; Koopmann, P.; Kovtunova, A.; Krüger, A.; Popovic, A.; and Siahaan, I. S. R. 2024. Explaining reasoning results for OWL ontologies with Eevee. In Marquis, P.; Ortiz, M.; and Pagnucco, M., eds., *Proceedings of the 21st International Conference on Principles of Knowledge Representation and Reasoning, KR 2024, Hanoi, Vietnam, November 2-8, 2024*.
- Baader, F.; Horrocks, I.; Lutz, C.; and Sattler, U. 2017. *An Introduction to Description Logic*. Cambridge University Press.
- Baader, F.; Koopmann, P.; Kriegel, F.; and Nuradiansyah, A. 2022. Optimal ABox repair w.r.t. static $\mathcal{EL}$ TBoxes: From quantified ABoxes back to ABoxes. In *The Semantic Web – 19th International Conference, ESWC 2022*, volume 13261 of *Lecture Notes in Computer Science*, 130–146. Springer.
- Bienvenu, M., and Bourgaux, C. 2016. Inconsistency-tolerant querying of description logic knowledge bases. In *Reasoning Web: Logical Foundation of Knowledge Graph Construction and Query Answering – 12th International Summer School 2016*, volume 9885 of *Lecture Notes in Computer Science*, 156–202. Springer.
- Bienvenu, M., and Maniere, Q. 2026. Data complexity of querying description logic knowledge bases under cost-based semantics. In *Proceedings of the Annual AAAI Conference on Artificial Intelligence*. Singapore, Singapore: AAAI Press.
- Bienvenu, M., and Rosati, R. 2013. Tractable approximations of consistent query answering for robust ontology-based data access. In Rossi, F., ed., *IJCAI 2013, Proceedings of the 23rd International Joint Conference on Artificial Intelligence, Beijing, China, August 3-9, 2013*, 775–781. IJCAI/AAAI.
- Bienvenu, M.; Bourgaux, C.; and Goasdoué, F. 2019. Computing and explaining query answers over inconsistent DL-Lite knowledge bases. *Journal of Artificial Intelligence Research* 64:563–644.
- Bienvenu, M.; Inoue, K.; and Kozhemiachenko, D. 2024. Abductive reasoning in a paraconsistent framework. In Marquis, P.; Ortiz, M.; and Pagnucco, M., eds., *Proceedings of the 21st International Conference on Principles of Knowledge Representation and Reasoning, KR 2024*.
- Calvanese, D.; Ortiz, M.; Simkus, M.; and Stefanoni, G. 2013. Reasoning about explanations for negative query answers in DL-Lite. *J. Artif. Intell. Res.* 48:635–669.
- Ceylan, İ. İ.; Lukasiewicz, T.; Malizia, E.; Molinaro, C.; and Vaisentavicius, A. 2020. Explanations for negative query answers under existential rules. In *Proceedings of the 17th International Conference on Principles of Knowledge Representation and Reasoning, KR 2020*, 223–232. AAAI Press.
- Du, J.; Wan, H.; and Ma, H. 2017. Practical TBox abduction based on justification patterns. In *Proceedings of the Thirty-First AAAI Conference on Artificial Intelligence*, 1100–1106.
- Du, J.; Wang, K.; and Shen, Y. 2015. Towards tractable and practical ABox abduction over inconsistent description logic ontologies. In *Proceedings of the Twenty-Ninth AAAI Conference on Artificial Intelligence*, 1489–1495. AAAI Press.
- Elsenbroich, C.; Kutz, O.; and Sattler, U. 2006. A case for abductive reasoning over ontologies. In *Proceedings of the OWLED’06 Workshop on OWL: Experiences and Directions*, volume 216 of *CEUR Workshop Proceedings*. CEUR-WS.org.
- Haak, A.; Koopmann, P.; Mahmood, Y.; and Turhan, A. 2025. Why not? Developing ABox abduction beyond repairs. In Tendra, L.; Ibáñez-García, Y.; and Koopmann, P., eds., *Proceedings of the 38th International Workshop on Description Logics - DL 2025, Opole, Poland, September 3-6, 2025*, CEUR Workshop Proceedings. CEUR-WS.org.
- Haak, A.; Koopmann, P.; Mahmood, Y.; and Turhan, A. 2026. Abox abduction for inconsistent knowledge bases under repair semantics. *CoRR* abs/2605.01341.
- Haifani, F.; Koopmann, P.; Turret, S.; and Weidenbach, C. 2022. Connection-minimal abduction in $\mathcal{EL}$ via translation to FOL. In *Proceedings of the 11th International Joint Conference on Automated Reasoning IJCAR 2022*, volume 13385 of *Lecture Notes in Computer Science*, 188–207. Springer.

- Ignatiev, A.; Narodytska, N.; and Marques-Silva, J. 2019. Abduction-based explanations for machine learning models. In *Proceedings of the AAAI Conference on Artificial Intelligence*, volume 33, 1511–1519.
- Klarman, S.; Endriss, U.; and Schlobach, S. 2011. ABox abduction in the description logic $\mathcal{ALC}$. *Journal of Automated Reasoning* 46(1):43–80.
- Koopmann, P.; Del-Pinto, W.; Turrett, S.; and Schmidt, R. A. 2020. Signature-based abduction for expressive description logics. In *Proceedings of the 17th International Conference on Principles of Knowledge Representation and Reasoning, KR 2020*, 592–602. AAAI Press.
- Koopmann, P.; Mahmood, Y.; Ngomo, A.-C. N.; and Tiwari, B. 2026. Can you tell the difference? contrastive explanations for abox entailments. In *Proceedings of the AAAI Conference on Artificial Intelligence*, volume 40, 19189–19197.
- Koopmann, P. 2021. Signature-based abduction with fresh individuals and complex concepts for description logics. In *Proceedings of the Thirtieth International Joint Conference on Artificial Intelligence*, 1929–1935. ijcai.org.
- Lembo, D.; Lenzerini, M.; Rosati, R.; Ruzzi, M.; and Savo, D. F. 2015. Inconsistency-tolerant query answering in ontology-based data access. *J. Web Semant.* 33:3–29.
- Liberatore, P. 2005. Redundancy in logic I: CNF propositional formulae. *Artif. Intell.* 163(2):203–232.
- Lukasiewicz, T.; Malizia, E.; and Molinaro, C. 2022. Explanations for negative query answers under inconsistency-tolerant semantics. In *Proceedings of the Thirty-First International Joint Conference on Artificial Intelligence, IJCAI-22*, 2705–2711. ijcai.org.
- Maier, F.; Ma, Y.; and Hitzler, P. 2013. Paraconsistent OWL and related logics. *Semantic Web* 4(4):395–427.
- Peirce, C. 1878. Deduction, induction and hypothesis: Popular science monthly, v. 13.
- Schlobach, S., and Cornet, R. 2003. Non-standard reasoning services for the debugging of description logic terminologies. In *Proceedings of the Eighteenth International Joint Conference on Artificial Intelligence 2003*, 355–362. Morgan Kaufmann.
- Sipser, M. 1997. *Introduction to the theory of computation*. PWS Publishing Company.
- Wei-Kleiner, F.; Dragisic, Z.; and Lambrix, P. 2014. Abduction framework for repairing incomplete $\mathcal{EL}$ ontologies: Complexity results and algorithms. In *Proceedings of the Twenty-Eighth AAAI Conference on Artificial Intelligence*, 1120–1127. AAAI Press.