

A revision-based approach to handling inconsistency in description logics *

Guilin Qi^{1†}, Weiru Liu², David Bell²

¹ *Institute AIFB*

Universität Karlsruhe

D-76128 Karlsruhe, Germany

gqi@aifb.uni-karlsruhe.de

² *School of Electronics, Electrical Engineering and Computer Science*

Queen's University Belfast

Belfast, BT7 1NN, UK

{w.liu, da.bell}@qub.ac.uk

Abstract. Recently, the problem of inconsistency handling in description logics has attracted a lot of attention. Many approaches have been proposed to deal with this problem based on existing techniques for inconsistency management. In this paper, we first define two revision operators in description logics; one is called a weakening-based revision operator and the other is its refinement. Based on the revision operators, we then propose an algorithm to handle inconsistency in a *stratified* description logic knowledge base. We show that when the weakening-based revision operator is chosen, the resulting knowledge base of our algorithm is semantically equivalent to the knowledge base obtained by applying *refined conjunctive maxi-adjustment* (RCMA) which refines disjunctive maxi-adjustment (DMA), known to be a good strategy for inconsistency handling in classical logic.

Keywords: knowledge representation, semantic web, description logics, inconsistency handling, stratification

1. Introduction

Ontologies play a crucial role for the success of the Semantic Web (Berners-Lee, Hendler, and Lassila, 2001). Currently, a large number of ontologies have been developed in various research domains, or even within a single domain. Description logics (or DLs for short) comprise a family of knowledge representation languages with which to represent ontologies. In many cases, we may need to merge or integrate several ontologies into a single ontology, and even if the original ontologies are individually consistent, the merged ontology may be inconsistent. Current DL reasoners, such as RACER (Haarslev and Möller, 2005) and FaCT (Horrocks, 1998), can detect logical inconsistency. However,

* This paper was written when the first author was a PhD student at Queen's University Belfast.

† Corresponding author.



they only provide lists of unsatisfiable classes. The process of *resolving* inconsistency is left to the user or the ontology engineers. The need to improve DL reasoners to reason with inconsistency is becoming urgent to make them more applicable. Many approaches have been proposed to handle inconsistency in ontologies based on existing techniques for inconsistency management in traditional logics, such as propositional logic and nonmonotonic logics (Schlobach and Cornet, 2003; Parsia, Sirin, and Kalyanpur, 2005; Huang, Harmelen, and Teije, 2005).

It is well-known that priority or preference plays an important role in inconsistency handling (Baader and Hollunder, 1995b; Benferhat and Baida, 2004; Meyer, Lee, and Booth, 2005). In (Baader and Hollunder, 1995b), the authors introduced priority to default terminological logic such that more specific defaults are preferred to more general ones. When conflicts occur in reasoning with defaults, defaults which are more specific should be applied before more general ones. In (Meyer, Lee, and Booth, 2005), an algorithm, called *refined conjunctive maxi-adjustment* (RCMA for short) was proposed to weaken conflicting information in a *stratified* DL knowledge base and some consistent DL knowledge bases were obtained. To weaken a terminological axiom, they extend the DL $\mathcal{ALC}$ by a construct that can express global restrictions on the cardinality of concepts and weaken the axiom by relaxing the restrictions on the number of elements it may have. However, to weaken an assertional axiom, they simply delete it. An interesting problem is to seek and investigate other DL expressions to weaken a *conflicting* DL axiom (both terminological and assertional).

In this paper, we first define two revision operators in description logics: a weakening-based revision operator and its refinement. The revision operators are defined by *nominals*. The idea is that when a terminology axiom or a value restriction is in conflict, we simply add explicit exceptions to weaken it and assume that the number of exceptions is minimal. Based on the revision operators, we then propose an algorithm to handle inconsistency in a *stratified* description logic knowledge base. We show that when the weakening-based revision operator is chosen, the resulting knowledge base of our algorithm is semantically equivalent to that of the RCMA algorithm. However, their syntactical forms are different and they are applied to different description logics.

This paper is organized as follows. Section 2 gives a brief review of description logics. We then define two revision operators in Section 3. The revision-based algorithm for inconsistency handling is proposed in Section 4. Before the conclusion, we have a brief discussion on related work.

2. Description Logics

In this section, we introduce some basic notions of description logics (DLs), a well-known family of knowledge representation formalisms (Baader et al., 2003). DLs are fragments of first-order predicate logic. That is, they can be translated into first-order predicate logic (Borgida, 1994). They differ from their predecessors such as semantic networks and frames (Quillian, 1967; Minsky, 1981) in that they are equipped with a formal, logic-based semantics. In DLs, elementary descriptions are *concept names* (unary predicates) and *role names* (binary predicates). Complex descriptions are built from them inductively using concept and role constructors provided by the particular DL under consideration.

We consider the DL $\mathcal{ALC}$ (Schmidt-Schauß and Smolka, 1991) extended by nominals (Schaerf, 1994), which is a simple yet relatively expressive DL. $\mathcal{AL}$ is the abbreviation of attributive language and $\mathcal{C}$ denotes “complement”. Let N_C and N_R be pairwise disjoint and countably infinite sets of *concept names* and *role names* respectively. We use the letters A and B for concept names, the letter R for role names, and the letters C and D for concepts. $\top$ and $\perp$ denote the universal concept and the bottom concept respectively. The set of $\mathcal{ALC}$ concepts is the smallest set such that: (1) every concept name is a concept; (2) if C and D are concepts, R is a role name, then the following expressions are also concepts: $\neg C$ (full negation), $C \sqcap D$ (concept conjunction), $C \sqcup D$ (concept disjunction), $\forall R.C$ (value restriction on role names) and $\exists R.C$ (existential restriction on role names).

An interpretation $\mathcal{I} = (\Delta^{\mathcal{I}}, \cdot^{\mathcal{I}})$ consists of a set $\Delta^{\mathcal{I}}$, called the *domain* of $\mathcal{I}$, and a function $\cdot^{\mathcal{I}}$ which maps every concept C to a subset $C^{\mathcal{I}}$ of $\Delta^{\mathcal{I}}$ and every role R to a subset $R^{\mathcal{I}}$ of $\Delta^{\mathcal{I}} \times \Delta^{\mathcal{I}}$ such that, for all concepts C, D , role R , the following properties are satisfied:

- (1) $\top^{\mathcal{I}} = \Delta^{\mathcal{I}}$ and $\perp^{\mathcal{I}} = \emptyset$, $(\neg C)^{\mathcal{I}} = \Delta^{\mathcal{I}} \setminus C^{\mathcal{I}}$,
- (2) $(C \sqcap D)^{\mathcal{I}} = C^{\mathcal{I}} \cap D^{\mathcal{I}}$, $(C \sqcup D)^{\mathcal{I}} = C^{\mathcal{I}} \cup D^{\mathcal{I}}$,
- (3) $(\exists R.C)^{\mathcal{I}} = \{x \mid \exists y \text{ s.t. } (x, y) \in R^{\mathcal{I}} \text{ and } y \in C^{\mathcal{I}}\}$,
- (4) $(\forall R.C)^{\mathcal{I}} = \{x \mid \forall y (x, y) \in R^{\mathcal{I}} \text{ implies } y \in C^{\mathcal{I}}\}$.

For example, the concept description $Person \sqcap Female$ is an $\mathcal{ALC}$ -concept describing those persons that are female. Suppose $hasChild$ is a role name, the concept description $Person \sqcap \forall hasChild.Female$ expresses those persons whose children are all female. The concept $\forall hasChild.\perp \sqcap Person$ describes those persons who have no children.

A general concept inclusion axiom (GCI) or *terminology* is an inclusion statement of the form $C \sqsubseteq D$, where C and D are two (possibly complex) $\mathcal{ALC}$ concepts. It is the statement about how concepts are related to each other. An interpretation $\mathcal{I}$ satisfies a GCI $C \sqsubseteq D$ iff

$C^{\mathcal{I}} \subseteq D^{\mathcal{I}}$. A finite set of *GCI*s is called a *Tbox*. We can also formulate statements about individuals. We denote individual names as a, b, c . A *concept* (*role*) assertion axiom has the form $C(a)$ ($R(a, b)$), where C is a concept description, R is a role name, and a, b are *individual names*. An *Abox* contains a finite set of concepts and role axioms. In the Abox, one describes a specific state of affairs of an application domain in terms of concept and roles. To give a semantics to Aboxes, we need to extend interpretations to individual names. For each individual name a , $\cdot^{\mathcal{I}}$ maps it to an element $a^{\mathcal{I}} \in \Delta^{\mathcal{I}}$. The mapping $\cdot^{\mathcal{I}}$ should satisfy the *unique name assumption* (UNA)¹, that is, if a and b are distinct names, then $a^{\mathcal{I}} \neq b^{\mathcal{I}}$. An interpretation $\mathcal{I}$ satisfies a concept axiom $C(a)$ iff $a^{\mathcal{I}} \in C^{\mathcal{I}}$, it satisfies a role axiom $R(a, b)$ iff $(a^{\mathcal{I}}, b^{\mathcal{I}}) \in R^{\mathcal{I}}$. A DL knowledge base K consists of a Tbox and an Abox, i.e. it is a set of GCIs and assertion axioms. An interpretation $\mathcal{I}$ is a *model* of a DL (Tbox or Abox) axiom iff it satisfies this axiom, and it is a model of a DL knowledge base K if it satisfies every axiom in K . In the following, we use $M(\phi)$ (or $M(K)$) to denote the set of models of an axiom ϕ (or DL knowledge base K). K is consistent iff $M(K) \neq \emptyset$. Let K be an inconsistent DL knowledge base, a set $K' \subseteq K$ is a *conflict* of K if K' is inconsistent, and any sub-knowledge base $K'' \subset K'$ is consistent. Given a DL knowledge base K and a DL axiom ϕ , we say K *entails* ϕ , denoted as $K \models \phi$, iff $M(K) \subseteq M(\phi)$.

To define our approach, we need to extend $\mathcal{ALC}$ with nominals (Schaerf, 1994). A nominal has the form $\{a\}$, where a is an individual name. It can be viewed as a powerful generalization of DL Abox individuals. The nominal stands for exactly one individual. Semantically, it is different from an atomic concept, which is interpreted as some set of individuals. The semantics of $\{a\}$ is defined by $\{a\}^{\mathcal{I}} = \{a^{\mathcal{I}}\}$ for an interpretation $\mathcal{I}$. Nominals are included in many DLs, such as $\mathcal{SHOQ}$ (Horrocks and Sattler, 2001).

3. Revision Operators for DLs

3.1. DEFINITION

Belief revision is an important topic in knowledge representation. It deals with the problem of consistently accommodating new information received by an existing knowledge base. Recently, Flouris et al. discuss how to apply the famous AGM theory (Gardenfors, 1988) in

¹ In some very expressive DLs, such as $\mathcal{SHOQ}$, this assumption is dropped. Instead, they use *inequality assertions* of the form $a \neq b$ for individual names a and b , with the semantics that an interpretation $\mathcal{I}$ satisfies $a \neq b$ iff $a^{\mathcal{I}} \neq b^{\mathcal{I}}$.

belief revision to DLs and OWL (Flouris, Plexousakis and Antoniou, 2005). However, they evaluate only the feasibility of applying the *AGM postulates for contraction* in DLs. There is no explicit construction of a revision operator in their paper. In this section, we propose a revision operator for DLs and provide a semantic explanation of this operator.

We need some restrictions on the knowledge base to be revised. First, the original DL knowledge base should be consistent. This assumption is often accepted in belief revision theory. Second, we consider only inconsistencies arising due to objects explicitly introduced in the Abox. That is, suppose K and K' are the original knowledge base and the newly received knowledge base respectively, then for each conflict K_c of $K \cup K'$, K_c must contain an Abox statement. For example, we exclude the following case: $\top \sqsubseteq \exists R.C \in K$ and $\top \sqsubseteq \forall R.\neg C \in K'$. The handling of conflicting axioms in the Tbox has been discussed in some work recently (Schlobach and Cornet, 2003; Parsia, Sirin, and Kalyanpur, 2005). In this section, we discuss the resolution of conflicting information which contains assertional axioms in the context of knowledge revision.

We give a method to weaken a GCI first. To weaken a GCI, we simply add some explicit exceptions, and the number of exceptions is called the degree of the weakened GCI.

DEFINITION 1. *Let $C \sqsubseteq D$ be a GCI. A weakened GCI $(C \sqsubseteq D)_{weak}$ of $C \sqsubseteq D$ has the form $(C \sqcap \neg\{a_1\} \sqcap \dots \sqcap \neg\{a_n\}) \sqsubseteq D$, where n is the number of individuals to be removed from C . We use $d((C \sqsubseteq D)_{weak}) = n$ to denote the degree of $(C \sqsubseteq D)_{weak}$.*

It is clear that when $d((C \sqsubseteq D)_{weak}) = 0$, $(C \sqsubseteq D)_{weak} = C \sqsubseteq D$. The idea of weakening a GCI is similar to that of weakening an uncertain rule in (Benferhat and Baida, 2004). That is, when a GCI is involved in conflict, instead of dropping it completely, we remove those individuals which cause the conflict.

The weakening for an assertion is simpler than that for a GCI. The weakened assertion ϕ_{weak} of an Abox assertion $\phi = C(a)$ is of the form either $\phi_{weak} = \top(a)$ or $\phi_{weak} = \phi$. That is, we either weaken the concept C to $\top$ or keep it intact. The degree of ϕ_{weak} , denoted by $d(\phi_{weak})$, is defined as $d(\phi_{weak}) = 1$ if $\phi_{weak} = \top(a)$ and 0 otherwise.

Next, we consider the weakening of a DL knowledge base.

DEFINITION 2. *Let K and K' be two consistent DL knowledge bases. Suppose $K \cup K'$ is inconsistent. A DL knowledge base $K_{weak, K'}$ is a weakened knowledge base of K w.r.t K' if it satisfies:*

- $K_{weak, K'} \cup K'$ is consistent, and

- There is a bijection f from K to $K_{weak,K'}$ such that for each $\phi \in K$, $f(\phi)$ is a weakening of ϕ .

The set of all weakened bases of K w.r.t K' is denoted by $Weak_{K'}(K)$.

In Definition 2, the first condition requires that the weakened base should be consistent with K' . The second condition says that each element in $K_{weak,K'}$ is uniquely weakened from an element in K .

EXAMPLE 1. Let $K = \{bird(tweety), bird \sqsubseteq flies\}$ and $K' = \{\neg flies(tweety)\}$, where *bird* and *flies* are two concepts and *tweety* is an individual name. It is easy to check that $K \cup K'$ is inconsistent. Let $K_1 = \{\top(tweety), bird \sqsubseteq flies\}$, $K_2 = \{bird(tweety), bird \sqcap \neg\{tweety\} \sqsubseteq flies\}$, then both K_1 and K_2 are weakened bases of K w.r.t K' .

The degree of a weakened base is defined as the sum of the degrees of its elements.

DEFINITION 3. Let $K_{weak,K'}$ be a weakened base of a DL knowledge base K w.r.t K' . The degree of $K_{weak,K'}$ is defined as

$$d(K_{weak,K'}) = \sum_{\phi \in K_{weak,K'}} d(\phi)$$

In Example 1, we have $d(K_1) = d(K_2) = 1$.

We now define a revision operator.

DEFINITION 4. Let K be a consistent DL knowledge base. K' is a newly received DL knowledge base. The result of weakening-based revision of K w.r.t K' , denoted as $K \circ_w K'$, is defined as

$$K \circ_w K' = \{K' \cup K_i : K_i \in Weak_{K'}(K), \text{ and } \nexists K_j \in Weak_{K'}(K), d(K_j) < d(K_i)\}.$$

The result of revision of K by K' is a set of DL knowledge bases, each of which is the union of K' and a weakened base of K with the minimal degree. $K \circ_w K'$ is a *disjunctive DL knowledge base*² defined in (Meyer, Lee, and Booth, 2005).

We now consider the semantic aspect of our revision operator.

In (Meyer, Lee, and Booth, 2005), an ordering relation was defined to compare interpretations. It was claimed that only two interpretations having the same domain and mapping the same individual names to the same element in the domain can be compared. Given a domain

² A disjunctive DL knowledge base (or DKB) is a set of DL knowledge bases. A DKB $\mathcal{K}$ is satisfied by an interpretation $\mathcal{I}$ iff $\mathcal{I}$ is a model of at least one of the elements of $\mathcal{K}$.

Δ , a *denotation function* d is an injective mapping which maps every individual a to a different $a^{\mathcal{I}}$ in Δ . Then a *pre-interpretation* was defined as an ordered pair $\pi = (\Delta^{\pi}, d^{\pi})$, where Δ^{π} is a domain and d^{π} is a denotation function. For each interpretation $\mathcal{I} = (\Delta^{\mathcal{I}}, \cdot^{\mathcal{I}})$, its denotation function is denoted as $d^{\mathcal{I}}$. Given a pre-interpretation $\pi = (\Delta^{\pi}, d^{\pi})$, $\mathbf{I}^{\pi}$ is used to denote the class of interpretations $\mathcal{I}$ with $\Delta^{\mathcal{I}} = \Delta^{\pi}$ and $d^{\mathcal{I}} = d^{\pi}$. It is also assumed that a DL knowledge base is a multi-set³ of GCIs and assertion axioms. We now introduce the ordering between two interpretations defined in (Meyer, Lee, and Booth, 2005).

DEFINITION 5. *Let π be a pre-interpretation, $\mathcal{I} \in \mathbf{I}^{\pi}$, ϕ a DL axiom, and K a multi-set of DL axioms. If ϕ is an assertion, the number of ϕ -exceptions $e^{\phi}(\mathcal{I})$ is 0 if $\mathcal{I}$ satisfies ϕ and 1 otherwise. If ϕ is a GCI of the form $C \sqsubseteq D$, the number of ϕ -exceptions for $\mathcal{I}$ is:*

$$e^{\phi}(\mathcal{I}) = \begin{cases} |C^{\mathcal{I}} \cap (\neg D^{\mathcal{I}})| & \text{if } C^{\mathcal{I}} \cap (\neg D^{\mathcal{I}}) \text{ is finite} \\ \infty & \text{otherwise.} \end{cases} \quad (1)$$

The number of K -exceptions for $\mathcal{I}$ is $e^K(\mathcal{I}) = \sum_{\phi \in K} e^{\phi}(\mathcal{I})$. The ordering $\preceq_K^{\pi}$ on $\mathbf{I}^{\pi}$ is: $\mathcal{I} \preceq_K^{\pi} \mathcal{I}'$ iff $e^K(\mathcal{I}) \leq e^K(\mathcal{I}')$.

We give a proposition to show that our weakening-based revision operator captures some kind of minimal change.

PROPOSITION 1. *Let K be a consistent DL knowledge base. K' is a newly received DL knowledge base. Let Π be the class of all pre-interpretations. $\circ_w$ is the weakening-based revision operator. We then have*

$$M(K \circ_w K') = \cup_{\pi \in \Pi} \min(M(K'), \preceq_K^{\pi}).$$

Proposition 1 says that the models of the resulting knowledge base of our revision operator are models of K' which are minimal *w.r.t* the ordering $\preceq_K^{\pi}$ induced by K , i.e., models of K' with the minimal number of K -exceptions.

The proofs of proposition 1 and other propositions can be found in the appendix.

Let us illustrate the weakening-based revision operator by an example.

EXAMPLE 2. *Let $K = \{\forall \text{hasChild.RichHuman}(\text{Bob}), \text{hasChild}(\text{Bob}, \text{Mary}), \text{RichHuman}(\text{Mary}), \text{hasChild}(\text{Bob}, \text{Tom})\}$. Suppose we now receive new information $K' = \{\text{hasChild}(\text{Bob}, \text{John}), \neg \text{RichHuman}(\text{John})\}$. It is clear that $K \cup K'$ is inconsistent. Since $\forall \text{hasChild.Rich}$*

³ A multi-set is a set in which an element can appear more than once.

Human(Bob) is the only assertion axiom involved in conflict with K' , we only need to weaken it to restore consistency, that is, $K \circ_w K' = \{hasChild(Bob, Mary), RichHuman(Mary), hasChild(Bob, Tom), hasChild(Bob, John), \neg RichHuman(John), \top(Bob)\}$.

3.2. REFINED WEAKENING-BASED REVISION

In weakening-based revision, to weaken a conflicting assertion axiom, we weaken the concept to $\top$. However, this may result in counterintuitive conclusions. In Example 2, after revising K by K' using the weakening-based operator, we cannot infer that $RichHuman(Tom)$ because $\forall hasChild. RichHuman(Bob)$ is discarded, which is counterintuitive. From $hasChild(Bob, Tom)$ and $\forall hasChild. RichHuman(Bob)$ we should have known that $RichHuman(Tom)$ and this assertion is not in conflict with information in K' . The solution for this problem is to treat *John* as an *exception* and that all children of *Bob* other than *John* are rich humans.

Next, we propose a new method for weakening Abox assertions. For an Abox assertion of the form $\forall R.C(a)$, it is weakened by dropping some individuals which are related to the individual a by the relation R , i.e. its weakening has the form $\forall R.(C \sqcup \{b_1, \dots, b_n\})(a)$, where b_i ($i = 1, n$) are individuals to be dropped. For other Abox assertions ϕ , we either keep them intact or replace them by $\top(a)$.

DEFINITION 6. *Let ϕ be an assertion in an Abox. A weakened assertion ϕ_{weak} of ϕ is defined as:*

$$\phi_{weak} = \begin{cases} \forall R.(C \sqcup \{b_1, \dots, b_n\})(a) & \text{if } \phi = \forall R.C(a) \\ \top(a) \text{ or } \phi & \text{otherwise.} \end{cases} \quad (2)$$

The degree of ϕ_{weak} is $d(\phi_{weak}) = n$ if $\phi = \forall R.C$ and $\phi_{weak} = \forall R.(C \sqcup \{b_1, \dots, b_n\})(a)$, $d(\phi_{weak}) = 1$ if $\phi \neq \forall R.C$ and $\phi_{weak} = \top(a)$, and $d(\phi_{weak}) = 0$ otherwise.

We refer to the weakened base obtained by applying weakening of GCIs in Definition 1 and weakening of assertions in Definition 6 as a refined weakened base. We then replace the weakened base by the refined weakened base in Definition 4 and get a new revision operator, which we call a refined weakening-based revision operator and is denoted as

$\circ_{rw}$.

Let us have a look at Example 2 again.

EXAMPLE 3. *(Example 2 Continued) According to our discussion before, $\forall hasChild. RichHuman(Bob)$ is the only assertion axiom involved in conflict in K and *John* is the only exception which makes*

$\forall \text{hasChild.RichHuman}(\text{Bob})$ conflicting, so $K \circ_{rw} K' = \{\forall \text{hasChild.}(\text{RichHuman} \sqcup \{\text{John}\})(\text{Bob}), \text{hasChild}(\text{Bob}, \text{Mary}), \text{RichHuman}(\text{Mary}), \text{hasChild}(\text{Bob}, \text{Tom}), \text{hasChild}(\text{Bob}, \text{John}), \neg \text{RichHuman}(\text{John})\}$. We then can infer that $\text{RichHuman}(\text{Tom})$ from $K \circ_{rw} K'$.

To give a semantic explanation of the refined weakening-based revision operator, we need to define a new ordering between interpretations.

DEFINITION 7. Let π be a pre-interpretation, $\mathcal{I} \in \mathbf{I}^\pi$, ϕ a DL axiom, and K be a multi-set of DL axioms. If ϕ is an assertion of the form $\forall R.C(a)$, the number of ϕ -exceptions for $\mathcal{I}$ is:

$$e_r^\phi(\mathcal{I}) = \begin{cases} |R^\mathcal{I}(a^\mathcal{I}) \cap (\neg C^\mathcal{I})| & \text{if } R^\mathcal{I}(a^\mathcal{I}) \cap (\neg C^\mathcal{I}) \text{ is finite} \\ \infty & \text{otherwise,} \end{cases} \quad (3)$$

where $R^\mathcal{I}(a^\mathcal{I}) = \{b \in \Delta^\mathcal{I} : (a^\mathcal{I}, b) \in R^\mathcal{I}\}$. If ϕ is an assertion which is not of the form $\forall R.C(a)$, the number of ϕ -exceptions $e_r^\phi(\mathcal{I})$ is 0 if $\mathcal{I}$ satisfies ϕ and 1 otherwise. If ϕ is a GCI of the form $C \sqsubseteq D$, the number of ϕ -exceptions for $\mathcal{I}$ is:

$$e_r^\phi(\mathcal{I}) = \begin{cases} |C^\mathcal{I} \cap (\neg D^\mathcal{I})| & \text{if } C^\mathcal{I} \cap (\neg D^\mathcal{I}) \text{ is finite} \\ \infty & \text{otherwise.} \end{cases} \quad (4)$$

The number of K -exceptions for $\mathcal{I}$ is $e_r^K(\mathcal{I}) = \sum_{\phi \in K} e_r^\phi(\mathcal{I})$. The refined ordering $\preceq_{r,K}^\pi$ on $\mathbf{I}^\pi$ is: $\mathcal{I} \preceq_{r,K}^\pi \mathcal{I}'$ iff $e_r^K(\mathcal{I}) \leq e_r^K(\mathcal{I}')$.

We have the following propositions for the refined weakening-based revision operator.

PROPOSITION 2. Let K be a consistent DL knowledge base, and let K' be a newly received DL knowledge base. Let Π be the class of all pre-interpretations, and let $\circ_{rw}$ be the weakening-based revision operator. We then have

$$M(K \circ_{rw} K') = \cup_{\pi \in \Pi} \min(M(K'), \preceq_{r,K}^\pi).$$

Proposition 2 says that the refined weakening-based operator can be accomplished with minimal change.

PROPOSITION 3. Let K be a consistent DL knowledge base, and let K' be a newly received DL knowledge base. We then have

$$K \circ_{rw} K' \models \phi, \forall \phi \in K \circ_w K'.$$

By Example 3, the converse of Proposition 3 is false. Thus, we have shown that the resulting knowledge base of the refined weakening-based

revision contains more original information than that of the weakening-based revision.

4. A Revision-based Algorithm

In this section, we define an algorithm for handling inconsistency in a stratified DL knowledge base, i.e., where each element of the base is assigned a rank, based on the weakening-based revision operator. More precisely, a stratified DL knowledge base is of the form $\Sigma = \{K_1, \dots, K_n\}$, where for each $i \in \{1, \dots, n\}$, K_i is a finite multi-set of DL sentences. Sentences in each stratum K_i have the same rank or reliability, while sentences contained in K_j such that $j > i$ are seen as less reliable. The stratification can be either given by experts or computed automatically (Haase and Völker, 2005; Ma, Qi, Hitzler, Lin, 2007).

4.1. REVISION-BASED ALGORITHM

We first need to generalize the (refined) weakening-based revision by allowing the newly received DL knowledge base to be a disjunctive DL knowledge base. That is, we have the following definition.

DEFINITION 8. *Let K be a consistent DL knowledge base. $\mathcal{K}'$ is a newly received disjunctive DL knowledge base. The result of (refined) weakening-based revision of K w.r.t $\mathcal{K}'$, denoted as $K \circ_w \mathcal{K}'$, is defined as*

$$K \circ_w \mathcal{K}' = \{K' \cup K_{weak, K'} : K' \in \mathcal{K}', K_{weak, K'} \in Weak_{K'}(K) \ \& \ \nexists K_i \in Weak_{K'}(K), d(K_i) < d(K_{weak, K'})\}.$$

Revision-based Algorithm (R-Algorithm)

Input: a stratified DL knowledge base $\Sigma = \{K_1, \dots, K_n\}$, a (refined) weakening-based revision operator $\circ$ (i.e. $\circ = \circ_w$ or $\circ_{rw}$), a new DL knowledge base K

Result: a disjunctive DL knowledge base $\mathcal{K}$

begin

$\mathcal{K} \leftarrow K_1 \circ K$;

for $i = 2$ to n **do**

$\mathcal{K} \leftarrow K_i \circ \mathcal{K}$;

return $\mathcal{K}$

end

The idea originates from the revision-based algorithms proposed in (Qi, Liu, and Bell, 2005). That is, we start by revising the set of sentences in the first stratum using the new DL knowledge base K , and the result of revision is a disjunctive knowledge base. We then revise the set of sentences in the second stratum using the disjunctive knowledge base obtained by the first step, and so on.

EXAMPLE 4. Let $\Sigma = \{K_1, K_2\}$ and $K = \{\top \sqsubseteq \top\}$, where $K_1 = \{W(t), \neg F(t), B(c)\}$ and $K_2 = \{B \sqsubseteq F, W \sqsubseteq B\}$ (W , F , B , t and c abbreviate *Wing*, *Flies*, *Bird*, *Tweety* and *Chirpy*). Let $\circ = \circ_w$ in R-Algorithm. Since K_1 is consistent, we have $\mathcal{K} = K_1 \circ_w K = \{K_1\}$. Since $K_1 \cup K_2$ is inconsistent, we need to weaken K_2 . Let $K'_2 = \{B \sqcap \neg\{t\} \sqsubseteq F, W \sqsubseteq B\}$ and $K''_2 = \{B \sqsubseteq F, W \sqcap \neg\{t\} \sqsubseteq B\}$, so $K'_2, K''_2 \in \text{Weak}(K_2)$ and $d(K'_2) = d(K''_2) = 1$. It is easy to check that $K'_2 \cup K_1$ and $K''_2 \cup K_1$ are both consistent and they are the only weakened bases of K_2 which are consistent with K_1 . So $K_2 \circ_w \mathcal{K} = \{K_1 \cup K'_2, K_1 \cup K''_2\} = \{\{W(t), \neg F(t), B(c), B \sqcap \neg\{t\} \sqsubseteq F, W \sqsubseteq B\}, \{W(t), \neg F(t), B(c), B \sqsubseteq F, W \sqcap \neg\{t\} \sqsubseteq B\}\}$. It is easy to check that $F(c)$ can be inferred from $K_2 \circ_w \mathcal{K}$.

Based on Proposition 3, it is easy to prove the following proposition.

PROPOSITION 4. Let $\Sigma = \{K_1, \dots, K_n\}$ be a stratified DL knowledge base and K be a DL knowledge base. Suppose $\mathcal{K}_1$ and $\mathcal{K}_2$ are disjunctive DL knowledge bases resulting from the R-Algorithm using the weakening-based operator and the refined weakening-based operator respectively. We then have, for each DL axiom ϕ , if $\mathcal{K}_1 \models \phi$ then $\mathcal{K}_2 \models \phi$.

Proposition 4 shows that the resulting knowledge base of R-Algorithm *w.r.t* the refined weakening-based operator contains more information than that of R-Algorithm *w.r.t* the weakening-based operator.

In the following we show that if the weakening-based revision operator is chosen, then our revision-based approach is equivalent to the refined conjunctive maxi-adjustment (RCMA) approach (Meyer, Lee, and Booth, 2005). The RCMA approach is defined in a model-theoretical way as follows (Meyer, Lee, and Booth, 2005).

DEFINITION 9. Let $\Sigma = \{K_1, \dots, K_n\}$ be a stratified DL knowledge base. Let Π be the class of all pre-interpretations. Let $\pi \in \Pi$, $\mathcal{I}, \mathcal{I}' \in \mathbf{I}^\pi$. The lexicographically combined preference ordering $\preceq_{lex}^\pi$ is defined as $\mathcal{I} \preceq_{lex}^\pi \mathcal{I}'$ iff $\forall j \in \{1, \dots, n\}$, $\mathcal{I} \preceq_{K_j}^\pi \mathcal{I}'$ or $\mathcal{I} \prec_{K_i}^\pi \mathcal{I}'$ for some $i < j$. Then the set of models of the consistent DL knowledge base extracted from Σ by means of $\preceq_{lex}^\pi$ is $\cup_{\pi \in \Pi} \min(\mathbf{I}^\pi, \preceq_{lex}^\pi)$.

The following proposition shows that our revision-based approach is equivalent to the RCMA approach when the weakening-based revision operator is chosen.

PROPOSITION 5. *Let $\Sigma = \{K_1, \dots, K_n\}$ be a stratified DL knowledge base and $K = \{\top \sqsubseteq \top\}$. Let $\mathcal{K}$ be the resulting DL knowledge base of *R-Algorithm*. We then have*

$$M(\mathcal{K}) = \cup_{\pi \in \Pi} \min(\mathbf{I}^\pi, \preceq_{lex}^\pi).$$

In (Meyer, Lee, and Booth, 2005), an algorithm was proposed to compute the RCMA approach in a syntactical way. The main difference between our algorithm and the RCMA algorithm is that the strategies for resolving terminological conflicts are different. The RCMA algorithm uses a preprocess to transform all the GCIs $C_i \sqsubseteq D_i$ to cardinality restrictions (Baader, Buchheit, and Hollander, 1996) of the form $\leq_0 C_i \sqcap \neg D_i$, i.e. the concepts $C_i \sqcap \neg D_i$ do not have any elements. Then those conflicting cardinality restrictions $\leq_0 C_i \sqcap D_i$ are weakened by relaxing the restrictions on the number of elements C may have, i.e. a weakening of $\leq_0 C_i \sqcap D_i$ is of the form $\leq_n C_i \sqcap D_i$ where $n > 1$. The resulting knowledge base contains cardinality restrictions and assertions and is no longer a DL knowledge base in a strict sense. By contrast, our algorithm weakens the GCIs by introducing *nominal* and role constructors. So the resulting DL knowledge base of our algorithm still contains GCIs and assertions.

5. Related Work

This work is closely related to the work on inconsistency handling in propositional and first-order knowledge bases in (Benferhat et al., 2004; Benferhat and Baida, 2004), the work on knowledge integration in DLs in (Meyer, Lee, and Booth, 2005) and the work on revision-based inconsistency handling approaches in (Qi, Liu, and Bell, 2005). In (Benferhat et al., 2004), a very powerful approach, called the disjunctive maxi-adjustment (DMA) approach, was proposed for weakening conflicting information in a stratified propositional knowledge base. The basic idea of the DMA approach is that starting from the information in the lowest stratum where formulae have the highest level of priority, when inconsistency is encountered in the knowledge base, it weakens the conflicting information in the higher strata iteratively. When applied to a first-order knowledge base directly, the DMA approach is not satisfactory because some important information is lost. A new approach was proposed in (Benferhat and Baida, 2004). For a first-order

formula, called an *uncertain rule*, with the form $\forall x P(x) \Rightarrow Q(x)$, when it is involved in a conflict in the knowledge base, instead of deleting it completely, the formula is weakened by dropping the instances of this formula that are responsible for the conflict. The idea of weakening GCIs in Definition 1 is similar to this idea. In (Meyer, Lee, and Booth, 2005), the authors proposed an algorithm for inconsistency handling by transforming every GCI in a DL knowledge base into a cardinality restriction, and a cardinality restriction responsible for a conflict is weakened by relaxing the restrictions on the number of elements it may have. So their strategy of weakening GCIs is different from ours. Furthermore, we proposed a refined revision operator which weakens not only the GCIs but also assertions of the form $\forall R.A(a)$. The idea of applying revision operators to deal with inconsistency in a stratified knowledge base was proposed in (Qi, Liu, and Bell, 2005). However, this work is only applicable in propositional stratified knowledge bases. The R-Algorithm is a successful application of the algorithm to DL knowledge bases.

There is much other work on inconsistency handling in DLs (Baader and Hollunder, 1995b; Baader and Hollunder, 1995a; Parsia, Sirin, and Kalyanpur, 2005; Quantz and Royer, 1992; Haase et. al., 2005; Schlobach, 2005; Schlobach and Cornet, 2003; Friedrich and Shchekotykhin, 2005; Flouris, Plexousakis and Antoniou, 2005; Huang, Harmelen, and Teije, 2005). In (Baader and Hollunder, 1995a; Baader and Hollunder, 1995b), Reiter's default logic (Reiter, 1980) is embedded into terminological representation formalisms, where conflicting information is treated using *exceptions*. To deal with conflicting default rules, each rule is instantiated using individuals appearing in an Abox and two existing methods are applied to compute all extensions. However, in practical applications, when there is a large number of individual names, it is not advisable to instantiate the default rules. Moreover, only conflicting default rules are dealt with and it is assumed that there is no error in the Abox. This assumption is dropped in our algorithm, that is, an assertion in an Abox may be weakened when it is involved in a conflict. Another work on handling conflicting defaults can be found in (Quantz and Royer, 1992). The authors proposed a preference semantics for defaults in terminological logics. As pointed out in (Meyer, Lee, and Booth, 2005), this method does not provide a weakening of the original knowledge base and the formal semantics is not cardinality-based. Furthermore, it is also assumed that there is no error in the Abox. In recent years, several methods have been proposed to debug erroneous terminologies and have them repaired when inconsistencies are detected (Schlobach and Cornet, 2003; Schlobach, 2005; Parsia, Sirin, and Kalyanpur, 2005; Friedrich and Shchekotykhin, 2005). A

general framework for reasoning with inconsistent ontologies based on *concept relevance* was proposed in (Huang, Harmelen, and Teije, 2005). The idea is to select from an inconsistent ontology some consistent sub-theories based on a *selection function*, which is defined on the syntactic or semantic relevance. Then standard reasoning on the selected sub-theories is applied to find *meaningful* answers. A problem with debugging of erroneous terminologies methods in (Schlobach and Cornet, 2003; Schlobach, 2005; Parsia, Sirin, and Kalyanpur, 2005; Friedrich and Shchekotykhin, 2005) and the reasoning method in (Huang, Harmelen, and Teije, 2005) is that both approaches delete terminologies in a DL knowledge base to obtain consistent subbases, thus the structure of DL language is not exploited.

6. Conclusions and Further Work

In this paper, we have proposed a revision-based algorithm for handling inconsistency in description logics. We mainly considered the following issues:

1. A weakening-based revision operator was defined in both syntactical and semantic ways. Since the weakening-based revision operator may result in counter-intuitive conclusions in some cases, we defined a refined version of this operator by introducing additional expressions in DLs.
2. A revision-based algorithm was presented to handle inconsistency in a stratified knowledge base. When the weakening-based revision operator is chosen, the resulting knowledge base of our algorithm is semantically equivalent to that of the RCMA algorithm. The main difference between these algorithms is that the strategies for resolving terminological information are different.

There are several problems worthy of further investigation. Our R-Algorithm is based on two particular revision operators. Clearly, if a normative definition of revision operators in DLs is provided, then the R-Algorithm can be easily extended. Unfortunately, such a definition does not exist as yet. As far as we know, the first attempt to deal with this problem can be found in (Flouris, Plexousakis and Antoniou, 2005). However, the authors studied only the feasibility of AGM's postulates for a *contraction* operator and their results are not very positive. That is, they have showed that in many important DLs, such as $\mathcal{SHOIN}(\mathcal{D})$ and $\mathcal{SHIQ}$, it is impossible to define a contraction operator that satisfies the AGM postulates.

Appendix

Proof of Proposition 1: Before proving Proposition 1, we need to prove two lemmas.

LEMMA 1. *Let K and K' be two consistent DL knowledge bases and $\mathcal{I}$ be an interpretation such that $\mathcal{I} \models K'$. Suppose $K \cup K'$ is inconsistent. Let $l = \min(d(K_{weak,K'}) : K_{weak,K'} \in Weak_{K'}(K), \mathcal{I} \models K_{weak,K'})$. Then $e^K(\mathcal{I}) = l$.*

Proof: We need only to prove that for each $K_{weak,K'} \in Weak_{K'}(K)$ such that $\mathcal{I} \models K_{weak,K'}$ and $d(K_{weak,K'}) = l$, $e^K(\mathcal{I}) = d(K_{weak,K'})$.

(1) Let $\phi \in K$ be an assertion axiom of the form $C(a)$. Suppose $e^\phi(\mathcal{I}) = 1$, then $\mathcal{I} \not\models \phi$. Since $\mathcal{I} \models K_{weak,K'}$, $\phi \notin K_{weak,K'}$. So $\phi_{weak} = \top(a)$ and then $d(\phi_{weak}) = 1$. Conversely, suppose $d(\phi_{weak}) = 1$, then $\phi_{weak} = \top(a)$. We must have $\mathcal{I} \not\models \phi$. Otherwise, let $K''_{weak,K'} = (K_{weak,K'} \setminus \{\top(a)\}) \cup \{\phi\}$. Since $\mathcal{I} \models \phi$, then $K''_{weak,K'}$ is consistent. It is clear $d(K''_{weak,K'}) < d(K_{weak,K'})$, which is a contradiction. So $\mathcal{I} \not\models \phi$, we then have $e^\phi(\mathcal{I}) = 1$. Thus, $e^\phi = 1$ iff $d(\phi) = 1$.

(2) Let $\phi = C \sqsubseteq D$ be a GCI axiom and $\phi_{weak} = (C \sqsubseteq D)_{weak} \in K_{weak,K'}$. Suppose $d(\phi_{weak}) = n$. That is, $\phi_{weak} = C \sqcap \neg\{a_1, \dots, a_n\} \sqsubseteq D$. Since $\mathcal{I} \models K_{weak,K'}$, $\mathcal{I} \models \phi_{weak}$. Moreover, for any other weakening ϕ'_{weak} of ϕ , if $d(\phi'_{weak}) < n$, then $\mathcal{I} \not\models \phi'_{weak}$ (because otherwise, we find another weakening $K'_{weak,K'} = (K_{weak,K'} \setminus \{\phi_{weak}\}) \cup \{\phi'_{weak}\}$ such that $d(K'_{weak,K'}) < d(K_{weak,K'})$ and $\mathcal{I} \models K'_{weak,K'}$). Since $\mathcal{I} \models \phi_{weak}$, $C^{\mathcal{I}} \setminus \{a_1^{\mathcal{I}}, \dots, a_n^{\mathcal{I}}\} \subseteq D^{\mathcal{I}}$. For each a_i , we must have $a_i \in C$ and $a_i \notin D$. Otherwise, we can delete such a_i and obtain $\phi'_{weak} = C \sqcap \neg\{a_1, \dots, a_{i-1}, a_{i+1}, \dots, a_n\} \sqsubseteq D$ such that $d(\phi'_{weak}) < d(\phi_{weak})$ and $\mathcal{I} \models \phi'_{weak}$, which is a contradiction. So $|C^{\mathcal{I}} \cap \neg D^{\mathcal{I}}| \leq n$. Since for each a_i , let $\phi'_{weak} = C \sqcap \neg\{a_1, \dots, a_{i-1}, a_{i+1}, \dots, a_n\} \sqsubseteq D$, then $\mathcal{I} \not\models \phi'_{weak}$, so $|C^{\mathcal{I}} \cap \neg D^{\mathcal{I}}| \geq n$. Therefore, we have $|C^{\mathcal{I}} \cap \neg D^{\mathcal{I}}| = n = d(\phi_{weak})$.

(1) and (2) together show that $e^K(\mathcal{I}) = l$.

LEMMA 2. *Let K and K' be two consistent knowledge bases and $\mathcal{I}$ be an interpretation such that $\mathcal{I} \models K'$. Suppose $K \cup K'$ is inconsistent. Let $d_m = \min(d(K_{weak,K'}) : K_{weak,K'} \in Weak_{K'}(K))$. Then $\mathcal{I} \in \bigcup_{\pi \in \Pi} \min(M(K'), \preceq_K^\pi)$ iff $e^K(\mathcal{I}) = d_m$.*

Proof: “If Part”

Suppose $e^K(\mathcal{I}) = d_m$. By Lemma 1, for each $\mathcal{I}'$ such that $\mathcal{I}' \models K'$, $e^K(\mathcal{I}') = l$, where $l = \min(d(K_{weak,K'}) : K_{weak,K'} \in Weak_{K'}(K), \mathcal{I}' \models K_{weak,K'})$. That is, there exists $K_{weak,K'} \in Weak_{K'}(K)$ such that $\mathcal{I}' \models K_{weak,K'}$ and $e^K(\mathcal{I}') = d(K_{weak,K'})$. Since $d(K_{weak,K'}) \leq d_m$, we have $e^K(\mathcal{I}') \leq e^K(\mathcal{I})$. So $\mathcal{I} \in \bigcup_{\pi \in \Pi} \min(M(K'), \preceq_K^\pi)$.

“Only If Part”

Suppose $\mathcal{I} \in \bigcup_{\pi \in \Pi} \min(M(K'), \preceq_K^\pi)$. We need to prove that for all $\mathcal{I}' \models K'$, $e^K(\mathcal{I}) \leq e^K(\mathcal{I}')$. Suppose $\mathcal{I} \in \mathbf{I}^\pi$ for some $\pi = (\Delta^\pi, d^\pi)$. It is clear that $\forall \mathcal{I}' \in \mathbf{I}^\pi$, $e^K(\mathcal{I}) \leq e^K(\mathcal{I}')$. Now suppose $\mathcal{I}' \in \mathbf{I}^{\pi'}$ for some $\pi' \neq \pi$ such that $\pi' = (\Delta^{\pi'}, d^{\pi'})$. We further assume that $e^K(\mathcal{I}') = \min(e^K(\mathcal{I}_i) : \mathcal{I}_i \models K')$. Let $\text{Ind}(K)$ and $\text{Ind}(K')$ be sets of individual names appearing in K and K' respectively. By the unique name assumption, for each individual name a in $\text{Ind}(K) \cup \text{Ind}(K')$, there is a unique element a_1 in Δ^π and a unique element a_2 in $\Delta^{\pi'}$ such that $a^\pi = a_1$ and $a^{\pi'} = a_2$. For notational simplicity, we assume that $a^\pi = a^{\pi'} = a$ for every individual name a . So $\text{Ind}(K) \cup \text{Ind}(K') \subseteq \Delta^\pi \cap \Delta^{\pi'}$. We take an $\mathcal{I}'' \in \mathbf{I}^\pi$ which satisfies the following conditions: 1) for each concept C appearing in K , suppose $\Delta = C^\pi \cap (\text{Ind}(K) \cup \text{Ind}(K'))$, then $\Delta \subseteq C^{\mathcal{I}''}$; 2) $e^K(\mathcal{I}'') = \min(e^K(\mathcal{I}) : \mathcal{I} \models K' \ \mathcal{I} \in \mathbf{I}^\pi)$. We now prove $\Sigma_{\phi \in K} e^\phi(\mathcal{I}') = \Sigma_{\phi \in K} e^\phi(\mathcal{I}'')$. By 1) and 2), suppose ϕ is an assertion of the form $C(a)$, where C is a concept, then $a^{\mathcal{I}'} \in C^{\mathcal{I}'}$ iff $a^{\mathcal{I}''} \in C^{\mathcal{I}''}$, so $e^\phi(\mathcal{I}') = e^\phi(\mathcal{I}'')$. Suppose ϕ is a GCI of the form $C \sqsubseteq D$ and $b \in C^{\mathcal{I}'} \cap \neg D^{\mathcal{I}'}$. Then we must have $b \in \text{Ind}(K) \cup \text{Ind}(K')$. Otherwise, define $\mathcal{I}''' = (\Delta^{\mathcal{I}'} \setminus \{b\}, \cdot^{\mathcal{I}'''})$ such that for each concept name C , $C^{\mathcal{I}'''} = C^{\mathcal{I}'} \setminus \{b\}$ and for all R , $R^{\mathcal{I}'''} = R^{\mathcal{I}'} \setminus (\{(b, a_i) : a_i \in \Delta^{\mathcal{I}'}\} \cup \{(a_i, b) : a_i \in \Delta^{\mathcal{I}'}\})$. It is easy to check that $\mathcal{I}''' \models K'$ and $e^K(\mathcal{I}''') < e^K(\mathcal{I}')$, which is a contradiction. So $b \in C^{\mathcal{I}'} \cap \neg D^{\mathcal{I}'} \cap (\text{Ind}(K) \cup \text{Ind}(K'))$. Since $C^{\mathcal{I}'} \cap (\text{Ind}(K) \cup \text{Ind}(K')) = C^{\mathcal{I}''} \cap (\text{Ind}(K) \cup \text{Ind}(K'))$ and $D^{\mathcal{I}'} \cap (\text{Ind}(K) \cup \text{Ind}(K')) = D^{\mathcal{I}''} \cap (\text{Ind}(K) \cup \text{Ind}(K'))$, we have $C^{\mathcal{I}'} \cap \neg D^{\mathcal{I}'} \cap (\text{Ind}(K) \cup \text{Ind}(K')) = (C^{\mathcal{I}''} \cap \neg D^{\mathcal{I}''}) \cap (\text{Ind}(K) \cup \text{Ind}(K'))$. It follows that $b \in C^{\mathcal{I}''} \cap \neg D^{\mathcal{I}''} \cap (\text{Ind}(K) \cup \text{Ind}(K'))$. We then have $C^{\mathcal{I}'} \cap \neg D^{\mathcal{I}'} \subseteq C^{\mathcal{I}''} \cap \neg D^{\mathcal{I}''}$. Similarly, we can prove that $C^{\mathcal{I}''} \cap \neg D^{\mathcal{I}''} \subseteq C^{\mathcal{I}'} \cap \neg D^{\mathcal{I}'}$. So $C^{\mathcal{I}''} \cap \neg D^{\mathcal{I}''} = C^{\mathcal{I}'} \cap \neg D^{\mathcal{I}'}$. That is, $e^\phi(\mathcal{I}) = e^\phi(\mathcal{I}'')$. Thus, we can conclude that $e^K(\mathcal{I}') = e^K(\mathcal{I}'')$. Since $e^K(\mathcal{I}'') = e^K(\mathcal{I})$, we have $e^K(\mathcal{I}) = e^K(\mathcal{I}')$. Therefore, for all $\mathcal{I}' \models K'$, $e^K(\mathcal{I}) \leq e^K(\mathcal{I}')$. It is clear that there exists an $\mathcal{I}' \models K'$ such that $e^{\mathcal{I}'} = d_m$. So $e^K(\mathcal{I}) = d_m$.

We continue the proof of Proposition 1. Suppose $\mathcal{I} \models K \circ_w K'$, then $\mathcal{I} \models K' \cup K_{\text{weak}, K'}$, for some $K_{\text{weak}, K'} \in \text{Weak}_{K'}(K)$ such that $d(K_{\text{weak}, K'}) = d_m$ (d_m is defined in Lemma 2). By Lemma 1, $\mathcal{I} \models K'$ and $e^K(\mathcal{I}) = d_m$. By Lemma 2, $\mathcal{I} \in \bigcup_{\pi \in \Pi} \min(M(K'), \preceq_K^\pi)$. Conversely, suppose $\mathcal{I} \in \bigcup_{\pi \in \Pi} \min(M(K'), \preceq_K^\pi)$. By Lemma 2, $\mathcal{I} \models K'$ and $e^K(\mathcal{I}) = d_m$. By Lemma 1, $\mathcal{I} \models K' \cup K_{\text{weak}, K'}$, for some $K_{\text{weak}, K'} \in \text{Weak}_{K'}(K)$ such that $d(K_{\text{weak}, K'}) = d_m$. So $\mathcal{I} \models K \circ_w K'$. This completes the proof.

Proof of Proposition 2: The proof of Proposition 2 is similar to that of Proposition 1. The only problem is that we need to extend the proofs of Lemma 1 and Lemma 2 by considering the weakening of assertion axioms of the form $\forall R.C(a)$, which can be proved similar to the case of GCIs.

Proof of Proposition 3: We need only to prove that $M(K \circ_{rw} K') \subseteq M(K \circ_w K')$. Suppose $\mathcal{I} \models K \circ_{rw} K'$, then by Proposition 2, $\mathcal{I} \models K'$ and $e_r^K(\mathcal{I}) = \min(e_r^K(\mathcal{I}') : \mathcal{I}' \models K')$. We now prove that for any $\mathcal{I}' \neq \mathcal{I}$, $e_r^K(\mathcal{I}) \leq e_r^K(\mathcal{I}')$. Suppose ϕ is an assertion of the form $\forall R.C(a)$ and $e_r^\phi(\mathcal{I}) \geq 1$, then there exists b such that $b^\mathcal{I} \in R^\mathcal{I}(a^\mathcal{I}) \cap (\neg D^\mathcal{I})$. Since $\mathcal{I} \not\models \forall R.C(a)$, we have $e_r^\phi(\mathcal{I}) = 1$. Since $e_r^\phi(\mathcal{I}') \geq e_r^\phi(\mathcal{I})$, we have $e_r^\phi(\mathcal{I}') \geq 1$. Similarly, we have $e_r^\phi(\mathcal{I}') = 1$. So $e_r^\phi(\mathcal{I}) = e_r^\phi(\mathcal{I}')$. Suppose $e_r^\phi(\mathcal{I}) = 0$ and $e_r^\phi(\mathcal{I}') \geq 1$, then $e_r^\phi(\mathcal{I}) = 0 < 1 = e_r^\phi(\mathcal{I}')$. Thus, $e_r^\phi(\mathcal{I}) \leq e_r^\phi(\mathcal{I}')$. If ϕ is an assertion which is not of the form $\forall R.C(a)$ or it is a GCI, then it is easy to prove that $e_r^\phi(\mathcal{I}) = e_r^\phi(\mathcal{I}')$. Therefore, $e^K(\mathcal{I}) \leq e^K(\mathcal{I}')$. By Proposition 1, $\mathcal{I} \in M(K \circ_w K')$.

Proof of Proposition 5: Let $\mathbf{I}_1^\pi = \min(\mathbf{I}^\pi, \preceq_{K_1}^\pi)$, and $\mathbf{I}_i^\pi = \min(\mathbf{I}_{i-1}^\pi, \preceq_{K_i}^\pi)$ for all $i > 1$. It is clear that $M(\mathcal{K}) = \mathbf{I}_n^\pi$. So we only need to prove that $\mathbf{I}_n^\pi = \min(\mathbf{I}^\pi, \preceq_{lex}^\pi)$. Suppose $\mathcal{I} \in \mathbf{I}_n^\pi$, then we must have $\mathcal{I} \in \min(\mathbf{I}^\pi, \preceq_{lex}^\pi)$. Otherwise, there exists $\mathcal{I}' \in \mathbf{I}^\pi$ such that $\mathcal{I}' \prec_{lex} \mathcal{I}$. That is, there exists i such that $\mathcal{I}' \prec_{K_i}^\pi \mathcal{I}$ and $\mathcal{I}' \simeq_{K_j}^\pi \mathcal{I}$ for all $j < i$, where $\mathcal{I}' \simeq_{K_j}^\pi \mathcal{I}$ means $\mathcal{I}' \preceq_{K_j}^\pi \mathcal{I}$ and $\mathcal{I} \preceq_{K_j}^\pi \mathcal{I}'$. Since $\mathcal{I}' \prec_{K_i}^\pi \mathcal{I}$, it is clear that $\mathcal{I}, \mathcal{I}' \in \mathbf{I}_{i-1}^\pi$ by the definition of $\mathbf{I}_{i-1}^\pi$. Since $\mathcal{I} \in \mathbf{I}_n^\pi$, we have $\mathcal{I} \in \mathbf{I}_i^\pi = \min(\mathbf{I}_{i-1}^\pi, \preceq_{K_i}^\pi)$, which is contradictory to the assumption that $\mathcal{I}' \prec_{K_i}^\pi \mathcal{I}$. Thus we prove that $\mathbf{I}_n^\pi \subseteq \min(\mathbf{I}^\pi, \preceq_{lex}^\pi)$. Conversely, suppose $\mathcal{I} \in \min(\mathbf{I}^\pi, \preceq_{lex}^\pi)$, then we must have $\mathcal{I} \in \mathbf{I}_n^\pi$. Otherwise, there exists an i such that $\mathcal{I} \notin \mathbf{I}_i^\pi$ and $\mathcal{I} \in \mathbf{I}_j^\pi$ for all $j < i$. Suppose $\mathcal{I}' \in \mathbf{I}_i^\pi$, then $\mathcal{I}' \in \mathbf{I}_j^\pi$ for all $j < i$. We then have $\mathcal{I}' \simeq_{K_j}^\pi \mathcal{I}$ for all $j < i$. Since $\mathcal{I}' \in \mathbf{I}_i^\pi$ and $\mathcal{I} \notin \mathbf{I}_i^\pi$, it follows that $\mathcal{I}' \prec_{K_i}^\pi \mathcal{I}$. That is, $\mathcal{I}' \prec_{lex}^\pi \mathcal{I}$, which is a contradiction. Thus we prove that $\min(\mathbf{I}^\pi, \preceq_{lex}^\pi) \subseteq \mathbf{I}_n^\pi$. This completes the proof.

References

- Baader, F. and B. Hollunder. Embedding defaults into terminological knowledge representation formalisms, *Journal of Automated Reasoning*, 14(1):149-180, 1995.
- Baader, F. and B. Hollunder. Priorities on defaults with prerequisites, and their application in treating specificity in terminological default logic, *Journal of Automated Reasoning*, 15(1): 41-68, 1995.
- Baader, F., M. Buchheit and B. Hollunder. Cardinality restrictions on concepts. *Artificial Intelligence*, 88:195-213, 1996.
- Baader, F., D. Calvanese, D.L. McGuinness, D. Nardi and P. F. Patel-Schneider. *The Description Logic Handbook: Theory, implementation and application*, Cambridge University Press, 2003.
- Benferhat, S. and R.E. Baida. A stratified first order logic approach for access control. *International Journal of Intelligent Systems*, 19:817-836, 2004.
- Benferhat, S., S. Kaci, D.L. Berre and M.-A. Williams. Weakening conflicting information for iterated revision and knowledge integration. *Artificial Intelligence*, 153(1-2):339-371, 2004.

- Berners-Lee, T., J. Hendler and O. Lassila. The semantic web, *Scientific American*, 284(5):3443, 2001.
- Borgida, A. On the relationship between description logic and predicate logic, In *Proceedings of the 3th International Conference on Information and Knowledge Management (CIKM'94)*, 219-225, 1994.
- Flouris, G., D. Plexousakis and G. Antoniou. On applying the AGM theory to DLs and OWL, In *Proceedings of 4th International Conference on Semantic Web (ISWC'05)*, 216-231, 2005.
- Friedrich, G. and K.M. Shchekotykhin. A general diagnosis method for ontologies, In *Proceedings of 4th International Conference on Semantic Web (ISWC'05)*, 232-246, 2005.
- Gärdenfors, P. *Knowledge in Flux-Modeling the Dynamic of Epistemic States*, The MIT Press, Cambridge, Mass, 1988.
- Haarslev, V. and R. Möller. RACER system description, In *Proceedings of the First International Joint Conference on Automated Reasoning (IJCAR'01)*, 701-706, 2001.
- Haase, P., F. van Harmelen, Z. Huang, H. Stuckenschmidt and Y. Sure. A framework for handling inconsistency in changing ontologies, In *Proceedings of the 4th International Semantic Web Conference (ISWC'05)*, LNCA3729, 353-367, 2005.
- Haase, P. and J. Völker. Ontology Learning and Reasoning - Dealing with Uncertainty and Inconsistency, In *Proceedings of the International Workshop on Uncertainty Reasoning for the Semantic Web*, 45-55, 2005.
- Horrocks, I. The FaCT system, In de Swart, H., ed., *Tableaux'98, LNAI 1397*, 307-312, 1998.
- Horrocks, I. and U. Sattler. Ontology reasoning in the $\mathcal{SHOQ}(D)$ description logic, In *Proceedings of the 17th International Joint Conference on Artificial Intelligence (IJCAI'01)*, 199-204, 2001.
- Huang, Z., F. van Harmelen, and A. ten Teije. Reasoning with inconsistent ontologies, In *Proceedings of the 19th International Joint Conference on Artificial Intelligence (IJCAI'05)*, 254-259, 2005.
- Ma, Y., Qi, G., Hitzler, P., and Lin, Z. Measuring Inconsistency for Description Logics Based on Paraconsistent Semantics, In *European Conference on Symbolic and Quantitative Approaches to Reasoning with Uncertainty (ECSQARU'07)*, to appear, 2007.
- Meyer, T., K. Lee, and R. Booth. Knowledge integration for description logics, In *Proceedings of the 20th National Conference on Artificial Intelligence (AAAI'05)*, 645-650, 2005.
- Minsky, M. A framework for representing knowledge. In J. Haugeland (ed.) *Mind Design*. The MIT Press. 1981.
- Parsia, B., E. Sirin and A. Kalyanpur. Debugging OWL ontologies, In *Proceedings of the 14th International World Wide Web Conference (WWW'05)*, 633-640, 2005.
- Qi, G., W. Liu, and D.A. Bell. A revision-based approach to resolving conflicting information, In *Proceedings of the 21th Conference on Uncertainty in Artificial Intelligence (UAI'05)*, 477-484, 2005.
- Quantz, J. and V. Royer. A preference semantics for defaults in terminological logics, In *Proceedings of the 3th Conference on Principles of Knowledge Representation and Reasoning (KR'92)*, 294-305, 1992.
- Quillian, M.R. Word concepts: A theory and simulation of some basic capabilities. *Behavioral Science*, 12: 410-430, 1967.
- Reiter, R. A Logic for Default Reasoning, *Artificial Intelligence*, 13(1-2): 81-132, 1980.

- Schaerf, A. Reasoning with individuals in concept languages. *Data and Knowledge Engineering*, 13(2):141-176, 1994.
- Schlobach, S. and R. Cornet. Non-standard reasoning services for the debugging of description logic terminologies, In *Proceedings of the 18th International Joint Conference on Artificial Intelligence (IJCAI'2003)*, 355-360, 2003.
- Schlobach, S. Diagnosing terminologies, In *Proceedings of the 20th National Conference on Artificial Intelligence (AAAI'05)*, 670-675, 2005.
- Schmidt-Schauß, M. and G. Smolka. Attributive concept descriptions with complements, *Artificial Intelligence*, 48:1-26, 1991.