

A revision-based approach for handling inconsistency in description logics

Guilin Qi, Weiru Liu, David A. Bell

School of Electronics, Electrical Engineering and Computer Science
Queen's University Belfast
Belfast, BT7 1NN, UK
{G.Qi, 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 were 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 the weakening-based revision operator and the other is its refinement. The logical properties of the operators are analyzed. 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 the disjunctive maxi-adjustment (DMA), a good strategy for inconsistency handling in classical logic.

Introduction

Ontologies play a crucial role for the success of the Semantic Web (Berners-Lee, Hendler, and Lassila 2001). There are many representation languages for ontologies, such as description logics (or DLs for short) and F-logic (Staab and Studer 2004). Recently, the problem of inconsistency (or incoherence) handling in ontologies has attracted a lot of attention and research addressing this problem has been reported in many papers (Baader and Hollunder; Baader and Hollunder 1995; Parsia, Sirin, and Kalyanpur 2005; Haase et. al. 2005; Schlobach 2005; Schlobach and Cornet 2003; Flouris, Plexousakis and Antoniou 2005; Huang, Harmelen, and Teije 2005; Meyer, Lee, and Booth 2005; Friedrich and Shchekotykhin 2005). Inconsistency can occur due to several reasons, such as modelling errors, migration or merging ontologies, and ontology evolution. Current DL reasoners, such as RACER (Haarslev and Möller 2005) and FaCT (Horrocks 1998), can detect logical inconsistency. However, they only provide lists of unsatisfiable classes. The process of *resolving* inconsistency is left to the user or ontology engineers. The need to improve DL reasoners to reason with inconsistency is becoming urgent to make them more applicable. Many approaches were 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; Benferhat and Baida 2004; Meyer, Lee, and Booth 2005). In (Baader and Hollunder), 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 introduced a DL expression, called *cardinality restrictions* on concepts. However, to weaken an assertional axiom, they simply delete it. An interesting problem is to explore 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, one is called a weakening-based revision operator and the other is its refinement. The revision operators are defined by introducing a DL constructor called *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.

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 conclusion, we have a brief discussion on related work.

Description logics

In this section, we introduce some basic notions of Description Logics (DLs), a family of well-known knowledge rep-

resentation formalisms (Baader et al. 2003). To make our approach applicable to a family of interesting DLs, we consider the well-known DL $\mathcal{ALC}$ (Schmidt-Schauß and Smolka 1991), which is a simple yet relatively expressive DL. 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 concept. 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: $\neg C$, $C \sqcap D$, $C \sqcup D$, $\forall R.C$ and $\exists R.C$.

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) $(\neg C)^{\mathcal{I}} = \Delta^{\mathcal{I}} \setminus C^{\mathcal{I}}$,
- (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) $(\forall R.C)^{\mathcal{I}} = \{x \mid \forall y (x, y) \in R^{\mathcal{I}} \text{ implies } y \in C^{\mathcal{I}}\}$.

We introduce an extra expression of DLs called *nominals* (also called *individual names*) (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 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) and $\mathcal{SHOIQ}$ (Horrocks and Sattler 2005).

A general concept inclusion axiom (GCI) or *terminology* is of the form $C \sqsubseteq D$, where C and D are two (possibly complex) $\mathcal{ALC}$ concepts. An interpretation $\mathcal{I}$ satisfies a GCI $C \sqsubseteq D$ iff $C^{\mathcal{I}} \subseteq D^{\mathcal{I}}$. A finite set of GCIs 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. 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 $(a^{\mathcal{I}}, b^{\mathcal{I}}) \in R^{\mathcal{I}}$. An *Abox* contains a finite set of concept and role axioms. 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 con-

sistent. Given a DL knowledge base K and a DL axiom ϕ , we say K entails ϕ , denoted as $K \models \phi$, iff $M(K) \subseteq M(\phi)$.

Revision Operators for DLs

Definition

Belief revision is a very 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 (Gärdenfors 1988) in belief revision to DLs and OWL (Flouris, Plexousakis and Antoniou 2005). However, they only evaluate the feasibility of apply the *AGM postulates for contraction* in DLs. There is no explicit construction of a revision operator in their paper. In this subsection, 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. Second, we only consider 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 much 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 weaken 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 of an assertion is simpler than that of a GCI. The weakened assertion ϕ_{weak} of an Abox assertion ϕ is of the form either $\phi_{weak} = \top$ or $\phi_{weak} = \phi$. That is, we either delete it or keep it intact. The degree of ϕ_{weak} , denoted as $d(\phi_{weak})$, is defined as $d(\phi_{weak}) = 1$ if $\phi_{weak} = \top$ 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

¹In 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 $a^{\mathcal{I}} \neq b^{\mathcal{I}}$.

- 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 base 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 $K' = \{\top, bird \sqsubseteq flies\}$, $K'' = \{bird(tweety), bird \sqcap \neg\{tweety\} \sqsubseteq flies\}$, then both K' and K'' 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} is defined as

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

In Example 1, we have $d(K') = d(K'') = 1$.

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 base² 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 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 $\mathcal{I} = (\Delta^{\mathcal{I}}, \mathcal{I})$, 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 GCI and assertion axioms. We now introduce the ordering between two interpretations defined in (Meyer, Lee, and Booth 2005).

²A disjunctive DL knowledge (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 $\mathcal{K}$.

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

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 . The proofs of proposition 2 and other propositions can be found in the appendix.

Let us look at an example.

Example 2 Let $K = \{\forall hasChild.RichHuman(Bob), hasChild(Bob, Mary), RichHuman(Mary), hasChild(Bob, Tom)\}$. Suppose we now receive new information $K' = \{hasChild(Bob, John), \neg RichHuman(John)\}$. It is clear that $K \cup K'$ is inconsistent. Since $\forall hasChild.RichHuman(Bob)$ is the only assertion axiom involved in conflict with K' , we only need to delete it to restore consistency, that is, $K \circ_w K' = \{hasChild(Bob, Mary), RichHuman(Mary), hasChild(Bob, Tom), hasChild(Bob, John), \neg RichHuman(John)\}$.

Refined weakening-based revision

In weakening-based revision, to weaken a conflicting assertion axiom, we simply delete it. 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 $\top$.

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

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

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 $\phi_{weak} = \top$ and $d(\phi_{weak}) = 0$ otherwise.

We call 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 $\forall hasChild.RichHuman(Bob)$ conflicting, so $K \circ_{rw} K' = \{\forall hasChild.(RichHuman \sqcup \{John\})(Bob), hasChild(Bob, Mary), RichHuman(Mary), hasChild(Bob, Tom), hasChild(Bob, John), \neg RichHuman(John)\}$. We then can infer that $RichHuman(Tom)$ 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 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. K' is a newly received DL knowledge base. Let Π be the class of all pre-interpretations. $\circ_{rw}$ is 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. K' is 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 important information than that of the weakening-based revision.

Logical properties of the revision operators

In belief revision theory, a set of postulates or logical properties are proposed to characterize a “rational” revision operator. The most famous postulates are so-called AGM postulates (Gärdenfors 1988) which were reformulated in (Katsuno and Mendelzon 1992). We now generalize AGM postulates for revision to DLs.

Definition 8 Given two DL knowledge bases K and K' . A revision operator $\circ$ is said to be AGM-compliant if it satisfies the following properties:

- (R1) $K \circ K' \models \phi$ for all $\phi \in K'$
- (R2) If $K \cup K'$ is consistent, then $M(K \circ K') = M(K \cup K')$
- (R3) If K' is consistent, then $K \circ K'$ is also consistent
- (R4) If $M(K) = M(K_1)$ and $M(K') = M(K_2)$, then $M(K \circ K') = M(K_1 \circ K_2)$
- (R5) $M(K \circ K') \cap M(K'') \subseteq M(K \circ (K' \cup K''))$
- (R6) If $M(K \circ K') \cap M(K'')$ is not empty, then $M(K \circ (K' \cup K'')) \subseteq M(K \circ K') \cap M(K'')$

(R1) says that the new information must be accepted. (R2) requires that the result of revision be equivalent to the union of the existing knowledge base and the newly arrived knowledge base if this union is satisfiable. (R3) is devoted to the satisfiability of the result of revision. (R4) is the syntax-irrelevance condition. (R5) and (R6) together are used to ensure minimal change. (R4) states that the operator is independent of the syntactical form of both the original knowledge base and the new knowledge base. The following property is obviously weaker than (R4)

(R4') If $M(K_1) = M(K_2)$, then $M(K \circ K_1) = M(K \circ K_2)$.

Definition 9 A revision operator $\circ$ is said to be quasi-AGM compliant if it satisfies (R1)-(R3), (R4'), (R5-R6).

The following proposition tells us the logical properties of our revision operators.

Proposition 4 Given two DL knowledge bases K and K' . Both the weakening-based revision operator and the refined weakening-based revision operator are not AGM-compliant but they satisfy postulates (R1), (R2), (R3), (R4'), (R5) and (R6), that is, they are quasi-AGM compliant.

Proposition 4 is a positive result. Our revision operators satisfy all the AGM postulates except (R4), i.e. the syntax-irrelevant condition.

A Revision-based Algorithm

It is well-known that priorities or preferences play an important role in inconsistency handling (Baader and Hollunder; Benferhat and Baida 2004; Benferhat et al. 2004; Meyer,

Lee, and Booth 2005). In this section, we define an algorithm for handling inconsistency in a stratified DL knowledge base, i.e. 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 $\Sigma = K_1 \cup \dots \cup K_n$, 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.

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

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

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 $K = \{\top\}$, 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 $\mathcal{K} = K_1 \circ_w \{\top\} = \{K_1\}$. 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 $K'_2, K''_2 \in Weak(K_2)$ 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 5 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 R-Algorithm using the weakening-based operator and 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 5 shows that the resulting knowledge base of R-Algorithm w.r.t the refined weakening-based operator contains more important 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.

Definition 11 (Meyer, Lee, and Booth 2005) 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 6 Let $\Sigma = (K_1, \dots, K_n)$ be a stratified DL knowledge base and $K = \{\top\}$. 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 information 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 $\leq_0 C_i \sqcap \neg D_i$ are weakened by relaxing the restrictions on the number of elements C may have, i.e. a weakening of $\leq_0 C_i \sqcap \neg D_i$ is of the form $\leq_n C_i \sqcap \neg D_i$ 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.

Application to revising a stratified DL knowledge base

We can define two revision operators based on R-Algorithm. Let $\Sigma = (K_1, \dots, K_n)$ be a stratified knowledge base and

K be a new DL knowledge base. Let $\circ$ be the (refined) weakening-based revision operator. The prioritized (refined) weakening-based revision operator, denoted as $\circ^g$, is defined in a model-theoretic way as: $M(\Sigma \circ^g K) = \bigcup_{\pi \in \Pi} \min(\{\mathcal{I} \in \mathbf{I}^\pi : \mathcal{I} \models K\}, \preceq_{lex}^\pi)$. We now look at the logical properties of the newly defined operator.

Proposition 7 *Let Σ be a stratified DL knowledge base, K and K' be two DL knowledge bases. The revision operator $\circ^g$ satisfies the following properties:*

- (P1) *If K is satisfiable, then $\Sigma \circ^g K$ is satisfiable.*
- (P2) $\Sigma \circ^g K \models \phi$, for all $\phi \in K$.
- (P3) *If $M(\Sigma) \cap M(K)$ is not empty, then $M(\Sigma \circ^g K) = M(\Sigma) \cap M(K)$.*
- (P4) *Given a stratified DL knowledge base $K = \{S_1, \dots, S_n\}$, and two DL knowledge bases K and K' , if $K \equiv K'$, then $Mod(\Sigma \circ^g K) = Mod(\Sigma \circ^g K')$.*
- (P5) $M(\Sigma \circ^g K') \cap M(K'') \subseteq M(\Sigma \circ^g (K' \cup K''))$.
- (P6) *If $M(\Sigma \circ^g K') \cap M(K'')$ is not empty, then $M(\Sigma \circ^g (K' \cup K'')) \subseteq M(\Sigma \circ^g K') \cap M(K'')$.*

(P1)-(P3) correspond to Conditions (R1)-(R3) in Definition 8. (P4) is a generalization of the weakening condition (R4') of the principle of irrelevance of syntax. (P5) and (P6) are generalization of (R5) and (R6).

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 revising-based inconsistency handling approaches in (Qi, Liu, and Bell 2005). In (Benferhat et al. 2004), a very powerful approach, called 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 with the lowest stratum where formulae have highest level of priority, when inconsistency is encountered in the knowledge base, it weakens the conflicting information in those strata. 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 some of 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 not only weakens 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 are many other work on inconsistency handling in DLs (Baader and Hollunder; Baader and Hollunder 1995; Parsia, Sirin, and Kalyanpur 2005; Quantz and Royer 1992; Haase et. al. 2005; Schlobach 2005; Schlobach and Cornet 2003; Flouris, Plexousakis and Antoniou 2005; Huang, Harmelen, and Teije 2005; Friedrich and Shchekotykhin 2005). In (Baader and Hollunder 1995; Baader and Hollunder), Reiter's default logic (Reiter 1987) is embedded into terminological representation formalisms, where conflicting information is treated as *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 information in the Abox is absolutely true. 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 information in the Abox was absolutely true. 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.

Conclusions and Further Work

In this paper, we propose 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. The well-known AGM postulates are reformulated and we showed that our operators satisfy most of the postulates. Thus they have good logical properties.
3. 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 our algorithm and the RCMA algorithm is that the strategies for resolving terminological information are different.
4. Two revision operators were defined on stratified DL knowledge bases and their logical properties were analyzed.

There are many 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 R-Algorithm can be easily extended. Unfortunately, such a definition does not exist now. 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 only studied the feasibility of AGM's postulates for a contraction operator and their results are not so positive. That is, they 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. Moreover, they didn't apply AGM's postulates for a revision operator and explicit construction of a revision operator was not considered in their paper. We generalized AGM postulates for revision in Definition 8 and we showed that our operators satisfied most of the generalized postulates. An important future work is to construct a revision operator in DLs which satisfies all the generalized AGM postulates.

Proofs

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 only need 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. Suppose $e^\phi(\mathcal{I}) = 1$, then $\mathcal{I} \models \phi$. Since $\mathcal{I} \models K_{weak,K'}$, $\phi \notin K_{weak,K'}$. So $\phi_{weak} = \top$ and then $d(\phi_{weak}) = 1$. Conversely, suppose $d(\phi_{weak}) = 1$, then $\phi_{weak} = \top$. We must have $\mathcal{I} \models \phi$. Otherwise, let $K''_{weak,K'} = (K_{weak,K'} \setminus \{\top\}) \cup \{\phi\}$. 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 $\mathcal{I} \models \phi$, 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 $\phi'_{weak} = C \sqcap \{a_1, \dots, a_{i-1}, a_{i+1}, \dots, a_n\} \sqsubseteq D$ 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 $\phi'_{weak} = C \sqcap \{a_1, \dots, a_{i-1}, a_{i+1}, \dots, a_n\} \sqsubseteq D$, 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 $Ind(K)$ and $Ind(K')$ be sets of individual names appearing in K and K' respectively. By unique name assumption, for each individual name a in $Ind(K) \cup Ind(K')$, there is a unique element a_1 in $\Delta^{\mathcal{I}}$ and a unique element a_2 in $\Delta^{\mathcal{I}'}$ such that $a^{\mathcal{I}} = a_1$ and $a^{\mathcal{I}'} = a_2$. For notational simplicity, we assume that $a^{\mathcal{I}} = a^{\mathcal{I}'} = a$ for every individual name a . So $Ind(K) \cup Ind(K') \subseteq \Delta^\pi \cap \Delta^{\pi'}$. We take an $\mathcal{I}'' \in \mathbf{I}^\pi$ which satisfies the following conditions: 1) for each concept C appearing in K , suppose $\Delta = C^{\mathcal{I}} \cap (Ind(K) \cup Ind(K'))$, 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 $b \in Ind(K) \cup Ind(K')$. Otherwise, if we 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 $b \in C^{\mathcal{I}'} \cap \neg D^{\mathcal{I}'} \cap (Ind(K) \cup Ind(K'))$. Since $C^{\mathcal{I}'} \cap (Ind(K) \cup Ind(K')) = C^{\mathcal{I}''} \cap (Ind(K) \cup Ind(K'))$ and $D^{\mathcal{I}'} \cap (Ind(K) \cup Ind(K')) = D^{\mathcal{I}''} \cap (Ind(K) \cup Ind(K'))$, we have $C^{\mathcal{I}'} \cap \neg D^{\mathcal{I}'} \cap (Ind(K) \cup Ind(K')) = C^{\mathcal{I}''} \cap \neg D^{\mathcal{I}''} \cap (Ind(K) \cup Ind(K'))$. It follows that $b \in C^{\mathcal{I}''} \cap \neg D^{\mathcal{I}''} \cap (Ind(K) \cup Ind(K'))$. 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 $\mathcal{I} \models K' \cup K_{weak, K'}$, for some $K_{weak, K'} \in Weak_{K'}(K)$ such that $d(K_{weak, K'}) = d_m$ (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, $\mathcal{I} \models K' \cup K_{weak, K'}$, for some $K_{weak, K'} \in Weak_{K'}(K)$ such that $d(K_{weak, K'}) = d_m$. 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 only need 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 $\mathcal{I}' \neq \mathcal{I}$, $e^K(\mathcal{I}) \leq e^K(\mathcal{I}')$. Suppose ϕ is an assertion of the form $\forall R.C(a)$ and $e^\phi(\mathcal{I}) \geq 1$, 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 $e^\phi(\mathcal{I}) = 1$. Since $e_r^\phi(\mathcal{I}') \geq e_r^\phi(\mathcal{I})$, we have $e_r^\phi(\mathcal{I}') \geq 1$. Similarly, we have $e^\phi(\mathcal{I}') = 1$. So $e^\phi(\mathcal{I}) = e^\phi(\mathcal{I}')$. Suppose $e_r^\phi(\mathcal{I}) = 0$ and $e_r^\phi(\mathcal{I}') \geq 1$, then $e^\phi(\mathcal{I}) = 0 < 1 = e^\phi(\mathcal{I}')$. Thus, $e^\phi(\mathcal{I}) \leq e^\phi(\mathcal{I}')$. If ϕ is an assertion which is not of the form $\forall R.C(a)$ or a GCI, then it is easy to prove that $e^\phi(\mathcal{I}) = e^\phi(\mathcal{I}')$. Therefore, $e^K(\mathcal{I}) \leq e^K(\mathcal{I}')$. By Proposition 1, $\mathcal{I} \in M(K \circ_w K')$.

Proofs of Proposition 4 and Proposition 5: Proposition 4 and Proposition 5 are easily to be checked and we do not provide their proofs here.

Proof of Proposition 6: 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 $M(K) = \mathbf{I}_n^\pi$. 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 $\mathcal{I}' \prec_{lex}^\pi \mathcal{I}$. 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 $\mathcal{I}' \simeq_{K_j}^\pi \mathcal{I}$, 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

- F. Baader and B. Hollunder. Embedding defaults into terminological knowledge representation formalisms, *Journal of Automated Reasoning*, 14(1):149-180, 1995.
- F. Baader and B. Hollunder. Priorities on defaults with pre-requisites, and their application in treating specificity in terminological default logic, *Journal of Automated Reasoning*, 15(1): 41-68, 1995.
- F. Baader, M. Buchheit, and B. Hollander. Cardinality restrictions on concepts. *Artificial Intelligence*, 88:195-213, 1996.
- F. Baader, D.L. McGuinness, D. Nardi, and Peter Patel-Schneider. *The Description Logic Handbook: Theory, implementation and application*, Cambridge University Press, 2003.
- S. Benferhat, C. Cayrol, D. Dubois, L. Lang, and H. Prade. Inconsistency management and prioritized syntax-based entailment. In *Proceedings of IJCAI'93*, 640-645, 1993.
- S. Benferhat, and R.E. Baida. A stratified first order logic approach for access control. *International Journal of Intelligent Systems*, 19:817-836, 2004.
- S. Benferhat, S. Kaci, D.L. Berre, and M.A. Williams. Weakening conflicting information for iterated revision and knowledge integration. *Artificial Intelligence*, vol. 153(1-2):339-371, 2004.
- T. Berners-Lee, J. Hendler, and O. Lassila. The semantic web, *Scientific American*, 284(5):3443, 2001.
- G. Flouris, D. Plexousakis and G. Antoniou. On applying the AGM theory to DLs and OWL, In *Proc. of 4th International Conference on Semantic Web (ISWC'05)*, 216-231, 2005.
- G. Friedrich and K.M. Shchekotykhin. A General Diagnosis Method for Ontologies, In *Proc. of 4th International Conference on Semantic Web (ISWC'05)*, 232-246, 2005.
- P. Gärdenfors, *Knowledge in Flux-Modeling the Dynamic of Epistemic States*, The MIT Press, Cambridge, Mass, 1988.
- V. Haarslev and R. Möller, RACER System Description, In *IJCAR'01*, 701-706, 2001.
- P. Haase, F. van Harmelen, Z. Huang, H. Stuckenschmidt, and Y. Sure. A framework for handling inconsistency in changing ontologies, In *ISWC'05, LNCA3729*, 353-367, 2005.
- I. Horrocks. The FaCT system, In de Swart, H., ed., *Tableaux'98, LNAI 1397*, 307-312, 1998.

- I. Horrocks, and U. Sattler. Ontology reasoning in the $\mathcal{SHOQ}(D)$ description logic, In *Proceedings of IJCAI'01*, 199-204, 2001.
- I. Horrocks and U. Sattler. A tableaux decision procedure for $\mathcal{SHOIQ}$, In *Proc. of 19th International Joint Conference on Artificial Intelligence (IJCAI'05)*, 448-453, 2005.
- Z. Huang, F. van Harmelen, and A. ten Teije. Reasoning with inconsistent ontologies, In *Proceedings of IJCAI'05*, 254-259, 2005.
- H. Katsuno and A.O. Mendelzon. Propositional Knowledge Base Revision and Minimal Change, *Artificial Intelligence*, 52(3): 263-294, 1992.
- C. Lutz, C. Areces, I. Horrocks, and U. Sattler. Keys, nominals, and concrete domains, *Journal of Artificial Intelligence Research*, 23:667-726, 2005.
- T. Meyer, K. Lee, and R. Booth. Knowledge integration for description logics, In *Proceedings of AAAI'05*, 645-650, 2005.
- B. Nebel. *What is Hybrid in Hybrid Representation and Reasoning Systems?*, In F. Gardin and G. Mauri and M. G. Filippini, editors, *Computational Intelligence II: Proc. of the International Symposium Computational Intelligence 1989*, North-Holland, Amsterdam, 217-228, 1990.
- B. Parsia, E. Sirin and A. Kalyanpur. Debugging OWL ontologies, In *Proc. of WWW'05*, 633-640, 2005.
- J. Quantz and V. Royer. A Preference Semantics for Defaults in Terminological Logics, In *Proc. of the 3th Conference on Principles of Knowledge Representation and Reasoning (KR'92)*, 294-305, 1992.
- G. Qi, W. Liu, and D.A. Bell. A revision-based approach to resolving conflicting information, In *Proceedings of twenty-first Conference on Uncertainty in Artificial Intelligence (UAI'05)*, 477-484.
- R. Reiter. A Theory of Diagnosis from First Principles, *Artificial Intelligence*, 32(1): 57-95, 1987.
- A. Schaerf. Reasoning with individuals in concept languages. *Data and Knowledge Engineering*, 13(2):141-176, 1994.
- S. Schlobach, and R. Cornet. Non-standard reasoning services for the debugging of description logic terminologies, In *Proceedings of IJCAI'2003*, 355-360, 2003.
- S. Schlobach. Diagnosing Terminologies, In *Proc. of AAAI'05*, 670-675, 2005.
- M. Schmidt-Schauß, and G. Smolka. Attributive Concept descriptions with complements, *Artificial Intelligence*, 48:1-26, 1991.
- S. Staab and R. Studer. *Handbook on Ontologies*, International Handbooks on Information Systems, Springer, 2004.
- H. Wang, A.L. Rector, N. Drummond and J. Seidenberg. Debugging OWL-DL Ontologies: A Heuristic Approach, In *Proc. of 4th International Conference on Semantic Web (ISWC'05)*, 745-757, 2005.