

Justification Masking in OWL

Matthew Horridge¹, Bijan Parsia¹, Ulrike Sattler¹

The University of Manchester, UK

Abstract. This paper presents a discussion on the phenomena of masking in the context of justifications for entailments. Various types of masking are introduced and a definition for each type is given.

1 Introduction

Many open source and commercial ontology development tools such as Protégé-4, Swoop, The NeOn Toolkit and Top Braid Composer use justifications [5] as a kind of explanation for entailments in ontologies. A *justification* for an entailment, also known as a *MinA* [1, 2], or a *MUPS* [11] if specific to explaining why a class name is unsatisfiable, is a minimal subset of an ontology that is sufficient for the given entailment to hold. More precisely, a justification is taken to be a subset minimal set of axioms that supports an entailment. Justifications are popular in the OWL world and, as the widespread tooling support shows, have been used in preference to full blown proofs for explaining why an entailment follows from a set of axioms.

However, despite the popularity of justifications, they suffer from several problems. Some of these problems, namely issues arising from the potential superfluity of axioms in justifications, were highlighted in [3]. Specifically, while all of the axioms in a justification are needed to support the entailment in question, there may be *parts* of these axioms that are not required for the entailment to hold. For example, consider $\mathcal{J} = \{A \sqsubseteq \exists R.B, \text{Domain}(R, C), C \sqsubseteq D \sqcap E\}$ which entails $A \sqsubseteq D$. While $\mathcal{J}$ is a justification for $A \sqsubseteq D$, and *all axioms* are required to support this entailment, there are *parts* of these axioms that are superfluous as far as the entailment is concerned: In the first axiom the filler of the existential restriction is superfluous, in the third axiom the conjunct E is superfluous for the entailment.

An important phenomenon related to superfluity has become known as *justification masking*. Recalling that there may be several justifications for an entailment, which may but do not have to overlap, masking refers to the case where the number of justifications for an entailment does not reflect the number of reasons for that entailment. For example, consider $\mathcal{J} = \{A \sqsubseteq \exists R.C \sqcap \forall R.C, D \equiv \exists R.C\}$ which entails $A \sqsubseteq D$. Clearly, $\mathcal{J}$ is a justification for $A \sqsubseteq D$. It is also noticeable that there are superfluous parts in this justification. Moreover, there are two distinct reasons why $\mathcal{J} \models A \sqsubseteq D$, the first being $\{A \sqsubseteq \exists R.C, \exists R.C \sqsubseteq D\}$ and the second being $\{A \sqsubseteq \exists R.\top \sqcap \forall R.C, \exists R.C \sqsubseteq D\}$. The work presented in the paper describes how masking can occur within a justification, over a set of justifications, and over a set of justifications and axioms outside justifications. The main

problems identified with masking are (i) it can hamper understanding—not all reasons for an entailment may be salient to a person trying to understand the entailment, and (ii) it can hamper the design or choice of a repair plan—not all reasons for an entailment may be obvious, and if the plan consists of weakening and removing parts of axioms it may not actually result in a successful repair of the ontology in question.

In [3] *laconic* and *precise* justifications were presented as a tool for dealing with the problems of superfluity and masking. However, while the basic *intuitions* of masking were presented in [3], and it was shown that laconic justifications could be used as a tool for working with masking, only two types of masking were discussed. This paper presents a comprehensive analysis of the different types of masking, provides a characterisation of masking, and lays down definitions and an analysis for the various types of masking.

2 Preliminaries

The work presented in this paper focuses on OWL 2. OWL 2 [8] is the latest standard in ontology languages from the World Wide Web Consortium. An OWL 2 ontology roughly corresponds to a $\mathcal{SROIQ}(D)$ [4] knowledge base. For the purposes of this paper, an *ontology* is regarded as a finite set of $\mathcal{SROIQ}$ axioms $\{\alpha_0, \dots, \alpha_n\}$. An axiom is of the form of $C \sqsubseteq D$ or $C \equiv D$, where C and D are (possibly complex) concept descriptions, or $S \sqsubseteq R$ or $S \equiv R$ where S and R are (possibly inverse or complex) roles.

It should be noted that OWL contains a significant amount of syntactic sugar, such as $\text{DisjointClasses}(C, D)$, $\text{FunctionalObjectProperty}(R)$ or $\text{Domain}(R, C)$. However, these axioms can be represented using sub-class and sub-property axioms.

Justifications are a popular form of explanation in the OWL world. A justification for an entailment η in an ontology $\mathcal{O}$, such that $\mathcal{O} \models \eta$ is a minimal subset of that entails η .

Definition 1 (Justification). $\mathcal{J}$ is a justification for $\mathcal{O} \models \eta$ if $\mathcal{J} \subseteq \mathcal{O}$, $\mathcal{J} \models \eta$ and for all $\mathcal{J}' \subsetneq \mathcal{J}$ $\mathcal{J}' \not\models \eta$.

By a slight abuse of notation, the nomenclature used in this paper also refers to a minimally entailing set of axioms (that is not necessarily a subset of an ontology) as a justification.

Much of the work presented in the remainder of the paper uses the “well known” structural transformation — δ . This transformation takes a set of axioms and flattens out each axiom by introducing names for sub-concepts, transforming the axioms into an equi-satisfiable set of axioms. The structural transformation was first described in Plaisted and Greenbaum [10], with a version of the rewrite rules for description logics given in [9].

In what follows, $\mathcal{A}$ is the ABox of an ontology, $\mathcal{R}$ is the RoleBox, and $\mathcal{T}$ is the TBox. A is an atomic concept in the signature of $\mathcal{O}$, A_D and A'_D are fresh concept names that are not in the signature of $\mathcal{O}$. C_i and D are arbitrary

concepts, excluding $\top$, $\perp$ and literals of the form X or $\neg X$ where X is not in the signature of $\mathcal{O}$, $\mathbf{C}$ is a possibly empty disjunction of arbitrary concepts. $C \equiv D$ is syntactic sugar for $C \sqsubseteq D$ and $D \sqsubseteq C$, as is $=nR.D$ for $\geq nR.D \sqcap \leq nR.D$. Domain and range axioms are GCIs so that $Domain(R, C)$ means $\exists R.\top \sqsubseteq C$, and $Range(R, C)$ means $\top \sqsubseteq \forall R.C$. The negation normal form of D is $\text{nnf}(D)$. The structural transformation δ is defined as follows:

$$\begin{aligned}
\delta(\mathcal{O}) &:= \bigcup_{\alpha \in \mathcal{R} \cup \mathcal{A}} \delta(\alpha) \cup \bigcup_{C_1 \sqsubseteq C_2 \in \mathcal{T}} \delta(\top \sqsubseteq \text{nnf}(\neg C_1 \sqcup C_2)) \\
\delta(D(a)) &:= \delta(\top \sqsubseteq \neg\{a\} \sqcup \text{nnf}(D)) \\
\delta(\top \sqsubseteq \mathbf{C} \sqcup D) &:= \delta(\top \sqsubseteq A'_D \sqcup \mathbf{C}) \cup \bigcup_{i=1}^{i=n} \delta(A'_D \sqsubseteq D_i) \text{ for } D = \bigcap_{i=1}^{i=n} D_i \\
\delta(\top \sqsubseteq \mathbf{C} \sqcup \exists R.D) &:= \delta(\top \sqsubseteq A_D \sqcup \mathbf{C}) \cup \{A_D \sqsubseteq \exists R.A'_D\} \cup \delta(A'_D \sqsubseteq D) \\
\delta(\top \sqsubseteq \mathbf{C} \sqcup \forall R.D) &:= \delta(\top \sqsubseteq A_D \sqcup \mathbf{C}) \cup \{A_D \sqsubseteq \forall R.A'_D\} \cup \delta(A'_D \sqsubseteq D) \\
\delta(\top \sqsubseteq \mathbf{C} \sqcup \geq nR.D) &:= \delta(\top \sqsubseteq A_D \sqcup \mathbf{C}) \cup \{A_D \sqsubseteq \geq nR.A'_D\} \cup \delta(A'_D \sqsubseteq D) \\
\delta(\top \sqsubseteq \mathbf{C} \sqcup \leq nR.D) &:= \delta(\top \sqsubseteq A_D \sqcup \mathbf{C}) \cup \{A_D \sqsubseteq \leq nR.A'_D\} \cup \delta(A'_D \sqsubseteq D) \\
\delta(A'_D \sqsubseteq D) &:= A'_D \sqsubseteq D \text{ (If } D \text{ is of the form } A \text{ or } \neg A) \\
\delta(A'_D \sqsubseteq D) &:= \delta(\top \sqsubseteq \neg A'_D \sqcup D) \text{ (If } D \text{ is not of the form } A \text{ or } \neg A) \\
\delta(\beta) &:= \beta \text{ for any other axiom}
\end{aligned}$$

The transformation ensures that concept names that are in the signature of $\mathcal{O}$ only appear in axioms of the form $X \sqsubseteq A$ or $X \sqsubseteq \neg A$, where X is some concept name *not* occurring in the signature of $\mathcal{O}$. Note that the structural transformation does not use structure sharing. For example, given $\top \sqsubseteq C \sqcup \exists R.C$, two new names are introduced, one for each use of C , to give $\{\top \sqsubseteq X_0 \sqcup X_1, X_0 \sqsubseteq C, X_1 \sqsubseteq \exists R.X_2, X_2 \sqsubseteq C\}$. The preclusion of structure sharing ensures that the different positions of C are captured.

The definition of laconic justifications uses the notion of the length of an axiom. Length is defined as follows: For X, Y a pair of concepts or roles, A a concept name, and R a role, the length of an axiom is defined as follows:

$$|X \sqsubseteq Y| := |X| + |Y|, \quad |X \equiv Y| := 2(|X| + |Y|),$$

where

$$\begin{aligned}
|\top| = |\perp| &:= 0, \\
|A| = |\{i\}| = |R| = |R^-| &:= 1, \\
|\neg C| &:= |C| \\
|C \sqcap D| = |C \sqcup D| &:= |C| + |D| \\
|\exists R.C| = |\forall R.C| = |\geq nR.C| = |\leq nR.C| &:= |R| + |C|
\end{aligned}$$

It should be noted that this definition is slightly different from the usual definition, but it allows cardinality axioms such as $A \sqsubseteq \leq 2R.C$ to be weakened to $A \sqsubseteq \leq 3R.C$ without increasing the length of the axiom.

In what follows the standard definition of deductive closure is used, and $\mathcal{O}^*$ is used to denote the deductive closure of $\mathcal{O}$.

Definition 2. $\mathcal{J}$ is a laconic justification for η over $\mathcal{O}$ if:

1. $\mathcal{J}$ is a justification in $\mathcal{O}^*$.
2. $\delta(\mathcal{J})$ is a justification in $(\delta(\mathcal{O}))^*$
3. For each $\alpha \in \delta(\mathcal{J})$ there is no α' such that
 - (a) α' is weaker than α ($\alpha \models \alpha'$ but $\alpha' \not\models \alpha$)
 - (b) $|\alpha'| \leq |\alpha|$
 - (c) $(\delta(\mathcal{J}) \setminus \{\alpha\}) \cup \delta(\alpha')$ is a justification for η

Intuitively, a laconic justification is a justification whose axioms do not contain any superfluous parts and all of whose parts are as weak as possible.

3 Intuitions about Masking

The basic notion of masking is that when taken on their own, the weakest parts of axioms in a justification may combine together with other parts of axioms within the justification or external to the justification to reveal further reasons that are not directly represented by the set of regular justifications, and do not directly have a one-to-one “correspondence” with the set of regular justifications.

We define four important types of masking: *Internal Masking*, *Cross Masking*, *External Masking* and *Shared Cores*. The intuitions behind these types of masking are explained below.

Internal Masking Internal masking refers to the phenomena where there are multiple reasons within a single justification as to why the entailment in question holds. An example of internal masking is shown below.

$$\mathcal{O} = \{A \sqsubseteq B \sqcap \neg B \sqcap C \sqcap \neg C\} \models A \sqsubseteq \perp$$

There is a single regular justification for $\mathcal{O} \models A \sqsubseteq \perp$, namely $\mathcal{O}$ itself. However, within this justification there are, intuitively, two reasons as to why $\mathcal{O} \models A \sqsubseteq \perp$, the first being $\{A \sqsubseteq B \sqcap \neg B\}$ and the second being $\{A \sqsubseteq C \sqcap \neg C\}$.

Cross Masking Intuitively, cross masking is present within a set of justifications for an entailment when parts of axioms from one justification combine with parts of axioms from another justification in the set to produce new reasons for the given entailment. For example, consider the following ontology.

$$\begin{aligned} \mathcal{O} = \{ & A \sqsubseteq B \sqcap \neg B \sqcap C \\ & A \sqsubseteq D \sqcap \neg D \sqcap \neg C \} \models A \sqsubseteq \perp \end{aligned}$$

There are two justifications for $\mathcal{O} \models A \sqsubseteq \perp$, namely $\mathcal{J}_1 = \{A \sqsubseteq B \sqcap \neg B \sqcap C\}$ and $\mathcal{J}_2 = \{A \sqsubseteq D \sqcap \neg D \sqcap \neg C\}$. However, part of the axiom in $\mathcal{J}_1$, namely $A \sqsubseteq C$ may combine with part of the axiom in $\mathcal{J}_2$, namely $A \sqsubseteq \neg C$ to produce a further reason: $\mathcal{J}_3 = \{A \sqsubseteq C, A \sqsubseteq \neg C\}$.

External Masking While internal masking and cross masking take place over a set of “regular” justifications for an entailment, external masking involves parts of axioms from a regular justification combining with parts of axioms from an ontology (intuitively the axioms outside of the set of regular justifications) to produce further reasons for the entailment in question. Consider the example below,

$$\begin{aligned} \mathcal{O} = \{ & A \sqsubseteq B \sqcap \neg B \sqcap C \\ & A \sqsubseteq \neg C \} \models A \sqsubseteq \perp \end{aligned}$$

There is just one justification for $\mathcal{O} \models A \sqsubseteq \perp$, however, although $A \sqsubseteq \neg C$ intuitively plays a part in the unsatisfiability of A it will never appear in a justification for $\mathcal{O} \models A \sqsubseteq \perp$. When $\mathcal{O}$ is taken into consideration, there are two salient reasons for $A \sqsubseteq \perp$, the first being $\{A \sqsubseteq B \sqcap \neg B\}$ and the second being $\{A \sqsubseteq C, A \sqsubseteq \neg C\}$

Shared Core Masking Finally, two justifications share a *core* if after stripping away the superfluous parts of axioms in each justification the justifications are essentially structurally equal. Consider the example below,

$$\begin{aligned} \mathcal{O} = \{ & A \sqsubseteq B \sqcap \neg B \sqcap C \\ & A \sqsubseteq B \sqcap \neg B \} \models A \sqsubseteq \perp \end{aligned}$$

There are two justifications for $\mathcal{O} \models \eta$, $\mathcal{J}_1 = \{A \sqsubseteq B \sqcap \neg B \sqcap C\}$ and $\mathcal{J}_2 = \{A \sqsubseteq B \sqcap \neg B\}$. However, $\mathcal{J}_1$ can be reduced to the laconic justification $\{A \sqsubseteq B \sqcap \neg B\}$ (since C is irrelevant for the entailment), which is structurally equal to $\mathcal{J}_2$. With regular justifications, it appears that there are more reasons for the entailment, when in fact each justification boils down to the same reason.

3.1 Masking Due to Weakening

The above intuitions have been illustrated using simple propositional examples. However, it is important to realise that masking is not just concerned with boolean parts of axioms. *Weakest parts* of axioms must also be taken into consideration. For example, consider

$$\begin{aligned} \mathcal{O} = \{ & A \sqsubseteq \geq 2R.C \\ & A \sqsubseteq \geq 1R.D \\ & C \sqsubseteq \neg D \} \models A \sqsubseteq \geq 2R \end{aligned}$$

There is one regular justification for $\mathcal{O} \models A \sqsubseteq \geq 2R$ namely, $\mathcal{J}_1 = \{A \sqsubseteq \geq 2R.C\}$. However, there are intuitively two reasons for this entailment. The first is described by the justification obtained as a weakening of $\mathcal{J}_1$, and is $\mathcal{J}_2 = \{A \sqsubseteq \geq 2R\}$. The second is obtained by weakening the first axiom in $\mathcal{O}$ and combining it with the second and third axioms in $\mathcal{O}$ to give $\{A \sqsubseteq \geq 1R.C, A \sqsubseteq \geq 1R.D, C \sqsubseteq \neg D\}$.

Of course, masking due to weakening can occur in internal masking, cross masking, external masking and shared cores.

3.2 Summary on Intuitions

As can be seen from the above examples, the *basic idea* is that when the weakest parts of axioms in a justification, set of justifications or an ontology are taken into consideration, there can be multiple reasons for an entailment that are otherwise not exposed with regular justifications. These reasons take the form of laconic justifications—justifications whose axioms do not contain any superfluous parts and whose parts are as weak as possible. With internal masking, cross masking and external masking, there are more laconic justifications (by some measure) than there are regular justifications. With shared cores there are fewer laconic justifications (by some measure) than there are regular justifications.

3.3 Detecting Masking

Given the above link between masking, weakest parts of axioms and laconic justifications, it may seem fruitful to use laconic justifications as a mechanism for detecting masking. Specifically, it may seem like a good idea to count laconic justifications for the entailment in question. However, this is a flawed intuition and several problems prevent laconic justification counting being used *directly* as a masking detection mechanism. We begin by noting that there may be an *infinite* number of laconic justifications for an entailment.

Lemma 1 (Number of Laconic Justifications). *Let $\mathcal{S}$ be a set of $SROIQ$ axioms such that $\mathcal{S} \models \eta$. In general, there may be an infinite number of laconic justifications over $\mathcal{S}$ for $\mathcal{S} \models \eta$.*

Proof: Consider an ontology $\mathcal{O}$ such that $\mathcal{O} \models A \sqsubseteq \perp$. Since laconic justifications may be drawn from the deductive closure of an ontology it is possible to construct an infinite set of justifications for the unsatisfiability of A of the form $\{A \sqsubseteq \geq nR.\top, A \sqsubseteq \leq (n-1)R.\top\}$.

The Promiscuity of the Deductive Closure The first problem is that, in general, there can be an infinite number of laconic justifications for a given entailment (Lemma 1). The notion of counting the number of laconic justifications over a set of axioms and comparing this to the number of regular justifications over the same set of axioms is therefore useless when it comes to detecting and defining masking. Even if the logic used did not result in an infinite number of laconic justifications, the effects of splitting and syntactic equivalence could result in miscounting. For example, consider $\mathcal{J}_1 = \{A \sqsubseteq B \sqcap C, B \sqcap C \sqsubseteq D\}$, where $\mathcal{J}_1$ is in itself laconic, however another justification $\mathcal{J}_2 = \{A \sqsubseteq B, A \sqsubseteq C, B \sqcap C \sqsubseteq D\}$ can be obtained, which is also laconic. Clearly, masking is not present in $\mathcal{J}_1$, but there are more laconic justifications than there are regular justifications.

Preferred Laconic Justifications Another approach might be to count the number of *preferred laconic justifications*, which are laconic justifications that

are made up of axioms which come from a filter on the deductive closure of a set of axioms. The notion of preferred laconic justifications was introduced in [3], where a filter called $\mathcal{O}^+$ is used to compute justifications that bear a syntactic resemblance to the axioms from which they are derived. Unfortunately, this idea is sensitive to the definition of the filter. Different filters, for different applications, may give different answers and false positives. While a particular filter could be verified to behave correctly and perhaps be used as an optimisation for detecting masking in an implementation, this mechanism is not appropriate for defining masking.

Preservation of Positional Information Another problem is that structural information can be lost with laconic justifications. Consider the $\{A \sqsubseteq B \sqcap (C \sqcap B)\}$ as a justification for $A \sqsubseteq B$. Masking is clearly present within this justification. If $B_{@1}$ denotes the first occurrence of B , and $B_{@2}$ denotes the second occurrence of B then A is a subclass of B because of two reasons: $A \sqsubseteq B_{@1}$ and $A \sqsubseteq B_{@2}$. However, this positional information is lost in all laconic justifications for $A \sqsubseteq B$. In essence, syntax is crucial when it comes to masking.

Splitting is Not Enough While syntax is very important when considering masking, it does not suffice to consider syntax alone. The example of masking due to weakening shows that simply splitting a set of axioms $\mathcal{S}$ into their constituent parts, using the structural transformation $\delta(\mathcal{S})$, and then examining the justifications for the entailment with respect $\delta(\mathcal{S})$ is not enough to capture this notion of masking. Weakenings of the split axioms must be considered in any mechanism that is used to detect masking.

4 Masking Defined

With the above intuitions and desiderata in mind the notion of masking can be made more concrete. The basic idea is to pull apart the axioms in a justification, set of justifications and an ontology, compute constrained weakenings of these parts (inline with the definition of laconic justifications), and then to check for the presence and number of laconic justifications within the set of regular justifications for an entailment with respect to these parts and their weakenings.

4.1 Parts and Their Weakenings

We first define a function $\delta^+(\mathcal{S})$, which maps a set of axioms $\mathcal{S}$ to a set of axioms composed from the union of $\delta(\mathcal{S})$ with the constrained weakenings of axioms in $\delta(\mathcal{S})$. The weakenings of axioms is constrained in that for an axiom $\alpha \in \delta(\mathcal{S})$, a weakening α' of α is contained in $\delta^+(\mathcal{S})$ only if α' is no longer than α —i.e. the weakening does not introduce any extra parts.

Definition 3 (δ^+). *For a set of $SRQIQ$ axioms, $\mathcal{S}$,*

$$\delta^+(\mathcal{S}) := \delta(\mathcal{S}) \cup \{\alpha' \mid \exists \alpha \in \delta(\mathcal{S}) \text{ s.t. } \alpha \models \alpha' \text{ and } \alpha' \not\models \alpha \text{ and } |\delta(\alpha')| = 1\}$$

Lemma 2 (δ^+ justificatory finiteness). *For a finite set of axioms $\mathcal{S}$, the set of justifications for an entailment in $\delta^+(\mathcal{S})$ is finite.*

Proof: $\delta^+(\mathcal{S})$ is composed of the set of axioms in $\delta(\mathcal{S})$, which is finite, plus a possibly infinite set of axioms taken from the deductive closure of *each axiom* in $\delta(\mathcal{S})$. For a *SRIOQ* axiom α , every axiom α' in $\delta(\alpha)$ must either be one of the following forms:

$$\begin{aligned} \top &\sqsubseteq X_i \sqcup X_j \\ X_i &\sqsubseteq A \\ X_i &\sqsubseteq \neg A \\ X_i &\sqsubseteq \exists R.X_j \\ X_i &\sqsubseteq \forall R.X_j \\ X_i &\sqsubseteq \{o\} \\ X_i &\sqsubseteq \exists R.\text{Self} \\ X_i &\sqsubseteq \geq nR.X_j \end{aligned}$$

in which case the set of axioms in $\delta^+(\alpha)$ is finite since the set of weakenings (in accordance with the definition of δ^+) of α' is finite. Or, α' is of the form:

$$X_i \sqsubseteq \leq nR.X_j$$

in which case there is an infinite number of weakenings of α' in $\delta^+(\alpha)$ since $A \sqsubseteq \leq (n+1)R.C$ is weaker than $A \sqsubseteq \leq nR.C$ for any $n \geq 0$. If justifications are made up solely of the axioms of the form corresponding to the first set then the set of justifications is clearly finite. If justifications contain axioms of the second form $X_i \sqsubseteq \leq nR.X_j$ then there is a finite upper bound m for n , where there are no justifications containing an axiom of the form $X_i \sqsubseteq \leq kR.X_j$ for some $k > m$. This is because, for values of k , where k is equal to the maximum number in $\leq$ restrictions in the closure of $\mathcal{S}$, or more, $X_i \sqsubseteq \leq kR.X_j$ is too weak to participate in a justification, and this follows as a straight forward consequence of *SRIOQ*'s model theory [4]. $\square$

Next, a function which filters out laconic justifications for an entailment from a set of justifications for the entailment is defined:

Definition 4 (Laconic Filtering). *For a set of axioms $\mathcal{S} \models \eta$, $\text{laconic}(\mathcal{S}, \eta)$ is the set of justifications for $\mathcal{S} \models \eta$ that are laconic over $\mathcal{S}$.*

Notice that because of Lemma 2, the set of justifications $\text{laconic}(\mathcal{S}, \eta)$ is finite.

4.2 Masking Definitions

With the definition of δ^+ and the definition of laconic filtering in hand, the various types of masking can now be defined.

Definition 5 (Internal Masking). For a justification $\mathcal{J}$ for $\mathcal{O} \models \eta$, internal masking is present within $\mathcal{J}$ if

$$|\text{laconic}(\delta^+(\mathcal{J}), \eta)| > 1$$

Lemma 3. Internal masking is not present within a laconic justification.

Proof: Assume that $\mathcal{J}$ is a laconic justification for η and that internal masking is present within $\mathcal{J}$. This means that there either must be (i) at least two laconic justifications for $\delta^+(\mathcal{J}) \models \eta$, i.e. there exists some $\mathcal{J}_1, \mathcal{J}_2 \subsetneq \delta^+(\mathcal{J})$ where $\mathcal{J}_1 \neq \mathcal{J}_2$ and are both laconic. However, since $\mathcal{J}$ itself is laconic this violates condition 2 of Definition 2, or (ii) there is a non-length increasing weakening of one or more axioms in $\delta(\mathcal{J})$ that yields $\delta(\mathcal{J})'$. However since $\mathcal{J}$ is laconic this violates conditions 3a and 3b of Definition 2. $\square$

Let $\mathcal{O} \models \eta$ and $\mathcal{J}_1, \dots, \mathcal{J}_n$ be the set of all justifications for $\mathcal{O} \models \eta$. Cross masking and External masking are then defined as follows:

Definition 6 (Cross Masking). For two justifications $\mathcal{J}_i$ and $\mathcal{J}_j$, cross masking is present within $\mathcal{J}_i$ and $\mathcal{J}_j$ if

$$|\text{laconic}(\delta^+(\mathcal{J}_i \cup \mathcal{J}_j), \eta)| > (|\text{laconic}(\delta^+(\mathcal{J}_i), \eta)| + |\text{laconic}(\delta^+(\mathcal{J}_j), \eta)|)$$

Definition 7 (External Masking). External masking is present if

$$|\text{laconic}(\delta^+(\mathcal{O}), \eta)| > |\text{laconic}(\delta^+(\bigcup_{i=1}^{i=n} \mathcal{J}_i), \eta)|$$

Definition 8 (Shared Cores). Two justifications $\mathcal{J}_i$ and $\mathcal{J}_j$ for $\mathcal{O} \models \eta$, share a core if there is a justification $\mathcal{J}'_i \in \text{laconic}(\delta^+(\mathcal{J}_i), \eta)$ and a justification $\mathcal{J}'_j \in \text{laconic}(\delta^+(\mathcal{J}_j), \eta)$ and a renaming ρ of terms not in $\mathcal{O}$ such that $\rho(\mathcal{J}'_i) = \mathcal{J}'_j$.

5 Examples

The issue of masking is indeed a real world problem with realistic ontologies. For example, external masking is present in the DOLCE ontology. The entailment quale $\sqsubseteq$ region has a single justification:

$$\{\text{quale} \equiv \text{region} \sqcap \exists \text{atomicPartOf.region}\}$$

However, there are further justifications that are externally masked by this regular justification. There are three laconic justifications, the first being

$$\{\text{quale} \sqsubseteq \text{region}\}$$

which is directly obtained as a weaker form of the regular justification. More interestingly, there are two additional laconic justifications:

$$\begin{aligned} &\{\text{quale} \sqsubseteq \exists \text{atomicPartOf.region} \\ &\text{atomicPartOf} \sqsubseteq \text{partOf} \\ &\text{partOf} \sqsubseteq \text{part}^- \\ &\text{region} \sqsubseteq \forall \text{part.region}\} \end{aligned}$$

and also

$$\begin{aligned} &\{\text{quale} \sqsubseteq \text{atomicPartOf.region} \\ &\text{atomicPartOf} \sqsubseteq \text{atomicPart}^- \\ &\text{atomicPart} \sqsubseteq \text{part} \\ &\text{region} \sqsubseteq \forall \text{part.region}\} \end{aligned}$$

Both of these justifications represent reasons for the entailment which are never seen with regular justifications due to the presence of external masking.

A real ontology about pathway interactions¹ contains an unsatisfiable class called “Phosphate Acceptor”. There are 32 regular justifications for this class being unsatisfiable. However, upon examination, these 32 justifications share a single core. When trying to understand the reason for the unsatisfiable class, the succinctness of the core provides a dramatic improvement in terms of usability.

6 Implementation Issues

The main focus of this paper has been to pin down the notions and types of masking. At this stage no attention has been paid to the practicalities of detection masking. However, the definitions for the various types of masking make use of the well known structural transformation $\delta \rightarrow \delta^+$ must be computed from δ . Naturally, this raises the question of performance and scalability, since many reasoners rely on the structure of axioms in real world ontologies for several key optimisations. Normalising the axioms in an ontology using the structural transformation, i.e. converting axioms to clausal form, raises the possibility of negating these optimisations. While more investigation work needs to be done, some preliminary experiments indicate that it is feasible to detect internal masking and cross masking. It is expected that an algorithm that transforms an ontology in an incremental manner, using techniques similar to those presented in [3] for computing laconic justifications, could provide a practical mechanism for detecting external masking.

7 Related Work

Various groups [6, 7, 11] have concentrated their efforts on what can be thought of as fine-grained justifications. In particular, Kalyanpur et al. [6, 5] presented work on fine-grained justifications, where axioms were split into smaller axioms in order to obtain a more “precise” justification. This work discusses the reasons for fine-grained justifications, one of which corresponds to the notion of external masking presented here. However, no precise definitions of masking were given in this work.

¹ <http://owl.cs.manchester.ac.uk/repository/download?ontology=http://purl.org/NET/biopax-obo/examples/reaction.owl> (courtesy Alan Ruttenberg)

8 Conclusions

This paper has presented a discussion on the phenomenon of justification masking. The notion and types of masking have been discussed and defined. These definitions basically identify the parts of axioms in a justification, over a set of justifications and an ontology, weaken the parts and then look for the number of laconic justifications that are present in the set of justifications over the axioms that represent these weakened parts.

References

1. Franz Baader and Bernhard Hollunder. Embedding defaults into terminological knowledge representation formalisms. *Journal of Automated Reasoning*, 14(1):149–180, 1995.
2. Franz Baader, Rafael Peñaloza, and Boontawee Suntisrivaraporn. Pinpointing in the description logic $\mathcal{EL}$. In Joachim Hertzberg, Michael Beetz, and Roman Engler, editors, *KI 2007: Advances in Artificial Intelligence, 30th Annual German Conference on AI, KI 2007, Osnabrück, Germany*, volume 4667 of *Lecture Notes in Computer Science*, pages 52–67. Springer, September 2007.
3. Matthew Horridge, Bijan Parsia, and Ulrike Sattler. Laconic and precise justifications in OWL. In *ISWC 08 The International Semantic Web Conference 2008, Karlsruhe, Germany*, 2008.
4. Ian Horrocks, Oliver Kutz, and Ulrike Sattler. The even more irresistible *SRQIQ*. In Patrick Doherty, John Mylopoulos, and Christopher A. Welty, editors, *The 10th International Conference on Principles of Knowledge Representation and Reasoning (KR 2006), Lake District, United Kingdom*, pages 57–67. AAAI Press, June 2006.
5. Aditya Kalyanpur. *Debugging and Repair of OWL Ontologies*. PhD thesis, The Graduate School of the University of Maryland, 2006.
6. Aditya Kalyanpur, Bijan Parsia, and Bernardo Cuenca Grau. Beyond asserted axioms: Fine-grain justifications for OWL-DL entailments. In Bijan Parsia, Ulrike Sattler, and David Toman, editors, *DL 2006, Lake District, U.K.*, volume 189 of *CEUR Workshop Proceedings*. CEUR-WS.org, 2006.
7. Joey Sik Chun Lam, Derek H. Sleeman, Jeff Z. Pan, and Wamberto Weber Vasconcelos. A fine-grained approach to resolving unsatisfiable ontologies. *Journal of Data Semantics*, 10:62–95, 2008.
8. Boris Motik, Peter F. Patel-Schneider, and Bijan Parsia. OWL 2 Web Ontology Language structural specification and functional style syntax. W3C Recommendation, W3C – World Wide Web Consortium, October 2009.
9. Boris Motik, Rob Shearer, and Ian Horrocks. Optimized reasoning in description logics using hypertableaux. In *Proc. of the 21st Int. Conf. on Automated Deduction (CADE-21)*, volume 4603 of *Lecture Notes in Artificial Intelligence*, pages 67–83. Springer, 2007.
10. David A. Plaisted and Steven Greenbaum. A structure-preserving clause form translation. *Journal of Symbolic Computation*, 1986.
11. Stefan Schlobach and Ronald Cornet. Non-standard reasoning services for the debugging of description logic terminologies. In *IJCAI International Joint Conference on Artificial Intelligence*, 2003.