

Minimizing Conservativity Violations in Ontology Alignments: Algorithms and Evaluation

Alessandro Solimando¹, Ernesto Jiménez-Ruiz², Giovanna Guerrini¹

¹DIBRIS, Informatica, Bioingegneria, Robotica e Ingegneria dei Sistemi; University of Genova; Italy

²Department of Computer Science; University of Oxford; United Kingdom

Abstract. In order to enable interoperability between ontology-based systems, ontology matching techniques have been proposed. However, when the generated mappings lead to undesired logical consequences, their usefulness may be diminished. In this paper, we present an approach to detect and minimize the violations of the so-called conservativity principle where novel subsumption entailments between named concepts in one of the input ontologies are considered as unwanted. The practical applicability of the proposed approach is experimentally demonstrated on the datasets from the Ontology Alignment Evaluation Initiative.

1. Introduction

Ontologies play a key role in the development of the Semantic Web and are being used in many diverse application domains, ranging from biomedicine to energy industry. An application domain may have been modeled with different points of view and purposes. This situation usually leads to the development of different ontologies that intuitively overlap, but they use different naming and modeling conventions.

The problem of (semi-)automatically computing mappings between independently developed ontologies is usually referred to as the *ontology matching problem*. A number of sophisticated ontology matching systems have been developed in the last years [16, 71]. Ontology matching systems, however, rely on lexical and structural heuristics and the integration of the input ontologies and the mappings may lead to many undesired logical consequences. In [36] three principles were proposed to minimize the number of potentially unintended consequences, namely: (i) *consistency principle*, the mappings should not lead to unsatisfiable concepts in the integrated ontology, (ii) *conservativity*

Received xxx

Revised xxx

Accepted xxx

principle, the mappings should not introduce new semantic relationships between concepts from one of the input ontologies, (iii) *locality principle*, the mappings should link entities that have similar *neighbourhoods*.

These alignment principles have been actively investigated in the last years (e.g., [32, 33, 36, 54, 56, 57, 66]). Violations to these principles are frequent, even in the reference mapping sets and the alignments generated by the best performing matchers of the Ontology Alignment Evaluation Initiative¹ (*OAEI*). Also manually curated alignments, such as the *UMLS Metathesaurus* [5] (*UMLS*),² a comprehensive effort for integrating biomedical knowledge bases, suffer from these violations [36]. The occurrence of these violations may hinder the usefulness of ontology mappings. The practical effect of these violations is clearly evident when ontology alignments are involved in complex tasks such as query answering [54, 78]. The undesired logical consequences caused by violations can either prevent query answering, or cause incorrect results. In order to reduce existing violations, alignment repair methods typically remove a subset of the alignment, given that input ontologies are considered as immutable, a common setting in ontology alignment repair scenarios.

It should be noted, however, the different nature of the alignment principles. Violations of the consistency principle, unlike violations of the conservativity and locality principles, always lead to an undesired logical consequence (i.e., unsatisfiability of a concept) and they should always be avoided. Conservativity and locality violations may also lead to undesired logical consequences; however they may also represent false positives and reveal incompleteness in one of the input ontologies. In Section 8 we discuss alternative approaches that suggest to fix the input ontologies instead of repairing the alignment (e.g., [7, 48]).

In this paper we focus on the conservativity violations and we follow a “better safe than sorry” approach (i.e., we treat violations as undesired consequences led by the mappings). Conservativity violations are presented in two flavours, namely *subsumption violations* and *equivalence violations*. The (potential) challenging number of conservativity violations requires to exploit the intrinsic characteristics of these two flavours, that result in the development of different approaches for their repair. The detection and correction of subsumption violations relies on the *assumption of disjointness* [67] and it is reduced to a consistency principle violation problem; while equivalence violations are addressed using a combination of graph theory and logic programming. These two methods are combined into a *multi-strategy* approach addressing both types of violations. Our extensive evaluation supports the effectiveness of the individual and combined approaches in the detection and correction of conservativity violations.

The present paper extends [74, 75] under the following aspects: all the experimental evaluations provided here cover both reference alignments and alignments computed by participating systems of the *OAEI* 2012–2014 campaigns, where previous papers covered only the reference alignments of the *OAEI*. Compared to [75], the present article fully details the proposed method, including a correctness proof of the technique for adding disjointness clauses to Horn Propositional formulas, on which our technique heavily relies. Furthermore, [75] only dealt with the subsumption violations flavour, while in this paper we also cover in detail the equivalence violations flavour. Concerning [74],³ all the technical details and proofs are now provided. In addition, the results of the evaluation of the two possible variants of our combined repair approach are now

¹ <http://oei.ontologymatching.org/>

² Alignments from UMLS are extracted according to the method defined in [36].

³ This paper was presented in a workshop without formal proceedings.

analyzed, as well as the results for the independent techniques in isolation, that can be used as baseline results. Finally, an empirical assessment of the impact of our repair methods on the alignment quality (in terms of precision, recall and f-measure) is now provided.

The remainder of the paper is organised as follows. Section 2 summarises the basic concepts and definitions we will rely on along the paper. In Section 3 we introduce our motivating scenario. Section 4 formally states the problem of computing repairs for equivalence violations and presents an algorithm to solve such violations. Section 5 describes the method and algorithm to solve subsumption violations. Section 6 details additional properties of the proposed methods. In Section 7 we present the conducted evaluation. A comparison with relevant related work is provided in Section 8. Finally, Section 9 gives some conclusions and future work lines.

2. Preliminaries

In this section, we provide the necessary definitions and notions that will be used in the subsequent sections. Section 2.1 briefly introduces OWL 2 and the main elements in an ontology. In Section 2.2 we give a formal definition of ontology mapping and ontology alignment (adapted from [17]) with their semantics. In Section 2.3 we precisely define the semantic consequences imposed by ontology alignments, and we formalize the *consistency* and *conservativity* principles. Finally, Section 2.4 covers the necessary preliminaries about graph theory.

2.1. Ontologies and OWL 2

Ontologies play a key role in the development of the Semantic Web and are being used in many diverse application domains, ranging from biomedicine to energy industry. The most widely used ontology modelling language is the OWL 2 Web Ontology Language [11], which is a World Wide Web Consortium (W3C) recommendation [84]. Description Logics (DL) are the formal underpinning of OWL 2 [3, 30].

An OWL 2 ontology $\mathcal{O}$ is equipped with a signature $Sig(\mathcal{O})$, that is a vocabulary of legal names for the *entities* appearing in the ontology. $Sig(\mathcal{O})$ is composed by the disjoint union of four finite sets: (i) N_C , a set of unary symbols called *named concepts*, (ii) N_R , a set of binary symbols called *named object properties*, (iii) N_D , a set of binary symbols called *data properties*, (iv) N_I , a set of constant symbols called *named individuals*.

OWL 2 ontologies can be seen as a set of axioms that are conformant to the syntactic rules and constraints imposed by their underlying DL language [30], and built using the elements of the signature. The *classification* of $\mathcal{O}$, denoted as $Cl(\mathcal{O})$, corresponds to the result of the computation, performed using an OWL 2 reasoner, of the full subsumption/subconcept relation between its named concepts (*i.e.*, elements of N_C). Classification is therefore the subset of the logical closure of an ontology $\mathcal{O}$ s.t. each axiom is of the form $A \sqsubseteq B$, where $A, B \in N_C(\mathcal{O})$ and $\mathcal{O} \models A \sqsubseteq B$.

2.2. Ontology Mappings and Alignments

Ontology Mappings. In Definition 2.1 we provide the definition of ontology *mapping* (also called *match* or *correspondence*).

Definition 2.1. Consider two input ontologies $\mathcal{O}_1, \mathcal{O}_2$, and their respective signature $Sig(\mathcal{O}_1)$ and $Sig(\mathcal{O}_2)$. A *mapping* between entities of $\mathcal{O}_1, \mathcal{O}_2$ is a 4-tuple $\langle e, e', r, c \rangle$ such that $e \in Sig(\mathcal{O}_1)$ and $e' \in Sig(\mathcal{O}_2)$, $r \in \{\sqsubseteq, \sqsupseteq, \equiv\}$ ⁴ is a semantic relation, and c is a confidence value. Usually, the real number unit interval $(0 \dots 1]$ is employed for representing confidence values. Mapping confidence intuitively reflects how reliable a mapping is (*i.e.*, 1 = very reliable, 0 = not reliable).

Ontology Alignment. Definition 2.2 introduces the notion of *alignment*.

Definition 2.2. An *alignment* $\mathcal{M}$ between two ontologies, namely $\mathcal{O}_1, \mathcal{O}_2$, is a set of mappings between $\mathcal{O}_1$ and $\mathcal{O}_2$.

The main format to represent mappings have been proposed in the context of the *Alignment API*, and it is called RDF Alignment [12]. This format is the standard for the well-known OAEI campaign. In addition, mappings are also represented as standard subclass and equivalence DL axioms. When mappings are expressed through OWL 2 axioms, confidence values are represented as OWL 2 axiom annotations [35]. The representation through standard OWL 2 axioms enables the reuse of the extensive range of OWL 2 reasoning infrastructure that is currently available. We adopt this representation, and in the remainder of the paper we consider alignments as set of OWL 2 axioms.

Definition 2.3 introduces the notion of *aligned ontology*, resulting from the integration of two input ontologies, through an alignment between them.

Definition 2.3. Let $\mathcal{O}_1, \mathcal{O}_2$ be two (input) ontologies, and let $\mathcal{M}$ be an alignment between them. The ontology $\mathcal{O}_{\mathcal{O}_1, \mathcal{O}_2}^{\mathcal{M}} = \mathcal{O}_1 \cup \mathcal{O}_2 \cup \mathcal{M}$ is called the *aligned ontology* w.r.t. $\mathcal{O}_1, \mathcal{O}_2$, and $\mathcal{M}$.

$\mathcal{O}_{\mathcal{O}_1, \mathcal{O}_2}^{\mathcal{M}}$ is simply called the aligned ontology when no confusion arises. Note that we assume that the signature of the aligned ontology is always the union of the signatures of the input ontologies. When the input ontologies are clear from the context we employ the abbreviated notation $\mathcal{O}^{\mathcal{M}}$.

Given that each mapping is translated into an OWL 2 axiom, the aligned ontology is again an OWL 2 ontology. Note that alternative formal semantics for ontology mappings have been proposed in the literature, such as those proposed by Zimmermann *et al.* in [87], and the semantics associated to the so-called *bridge rules*, in the context of distributed description logics [6, 55].

2.3. Semantics of the Integration and Principles for Ontology Alignments

This section introduces the semantics of the integration, and provides a formal characterization of the consistency and conservativity principles in ontology alignment.

Semantic Consequences of the Integration. The ontology resulting from the integration of two ontologies $\mathcal{O}_1$ and $\mathcal{O}_2$ via an alignment $\mathcal{M}$ may entail axioms that do not follow from $\mathcal{O}_1, \mathcal{O}_2$, or $\mathcal{M}$ alone. These new semantic consequences can be captured by the notion of *deductive difference* [46, 47].

Intuitively, the deductive difference between $\mathcal{O}$ and $\mathcal{O}'$, w.r.t. a signature Σ , is the set of entailments constructed over Σ that do not hold in $\mathcal{O}$, but do hold in $\mathcal{O}'$. The notion

⁴ We exclude disjointness from the semantic relations given that most of the available systems do not compute this relation. Negative constraints are typically harder to identify and assess than positive ones [20].

of deductive difference, however, has several drawbacks in practice. First, there is no algorithm available for computing it for DLs more expressive than $\mathcal{EL}$, for which the problem is already *EXPTIME*-complete [52]. In addition, the problem is undecidable for DLs as expressive as $\mathcal{ALCQIO}^5$ [51]. Second, the huge (possibly infinite) number of entailments in the difference is likely to overwhelm users.

In order to avoid the drawbacks of the deductive difference, we rely on the *approximation of the deductive difference*, specified in Definition 2.4.

Definition 2.4. Let A, B be named concepts (including $\top, \perp$), Σ be a signature, $\mathcal{O}$ and $\mathcal{O}'$ be two ontologies. We define the *approximation* of the Σ -deductive difference between $\mathcal{O}$ and $\mathcal{O}'$ (denoted $\text{diff}_{\Sigma}^{\approx}(\mathcal{O}, \mathcal{O}')$) as the set of axioms of the form $A \sqsubseteq B$ satisfying: (i) $A, B \in \Sigma$, (ii) $\mathcal{O} \not\models A \sqsubseteq B$, and (iii) $\mathcal{O}' \models A \sqsubseteq B$.

The proposed approximation only requires to compare the classification hierarchies of ontologies $\mathcal{O}$ and $\mathcal{O}'$ (i.e., $\text{Cl}(\mathcal{O})$ and $\text{Cl}(\mathcal{O}')$). In this paper we rely on this approximation, which has successfully been used in the past in the context of ontology integration to help users understanding the semantic consequences of this operation [35].

Consistency Principle Violations. Consistency violations evidence either a problem in the mappings or an incompatibility between the input ontologies, because they always result in incoherent and/or inconsistent (aligned) ontologies (i.e., an ontology containing unsatisfiable concepts). The consistency principle is the most widely investigated in the literature, where tools for detecting and automatically repair mappings leading to logical inconsistencies in the aligned ontology have been proposed (e.g., [33, 54]). Violations of the consistency principle are (typically) easy to detect with standard reasoning services. However, repairing such violations through standard reasoning services leads to intractability for medium size ontologies and alignments [34, 76, 77].

The consistency principle requires all the (named) concepts of the aligned ontology $\mathcal{O}_{\mathcal{O}_1, \mathcal{O}_2}^{\mathcal{M}}$ to be satisfiable, assuming the union of the input ontologies $\mathcal{O}_{\mathcal{O}_1, \mathcal{O}_2}^{\emptyset} = \mathcal{O}_1 \cup \mathcal{O}_2$ (without the alignment $\mathcal{M}$) does not contain unsatisfiable concepts. In Definition 2.5 we formally define (violations of) the consistency principle.

Definition 2.5. An alignment $\mathcal{M}$ violates the *consistency principle* (i.e., it is incoherent) w.r.t. $\mathcal{O}_1$ and $\mathcal{O}_2$ if $\text{diff}_{\Sigma}^{\approx}(\mathcal{O}_{\mathcal{O}_1, \mathcal{O}_2}^{\emptyset}, \mathcal{O}_{\mathcal{O}_1, \mathcal{O}_2}^{\mathcal{M}})$ contains axioms of the form $A \sqsubseteq \perp$, for all $A \in \Sigma = \text{Sig}(\mathcal{O}_1) \cup \text{Sig}(\mathcal{O}_2)$. Violations of the consistency principle result in an incoherent aligned ontology $\mathcal{O}_{\mathcal{O}_1, \mathcal{O}_2}^{\mathcal{M}}$.

Conservativity Principle Violations. The conservativity principle in ontology alignment aims at capturing the differences in the ontology classification between the input ontologies and the aligned ontology [36] (i.e., new subsumptions and/or new equivalences among concepts). The conservativity principle despite considering only ontology classification, and not the unrestricted problem addressed by conservative extensions [51], is of high interest because classification is one of the most used features in semantic-enabled applications [27, 49].

Conservativity violations may evidence, like consistency violations, erroneous mappings or disagreements in the input ontologies; however they may also reveal incompleteness in one of the ontologies. Although in the literature there are other approaches that consider these violations as false positives and suggest to fix the input ontologies

⁵ This DL is less expressive than $\mathcal{SROIQ}$, the underlying DL of OWL 2.

(see Section 8), in this paper we treat them as undesired consequences led by the mappings.

The conservativity principle (general notion) states that the aligned ontology $\mathcal{O}_{\mathcal{O}_1, \mathcal{O}_2}^{\mathcal{M}}$ should not induce any change in the concept hierarchies of the input ontologies $\mathcal{O}_1$ and $\mathcal{O}_2$. That is, the sets $\text{diff}_{\Sigma_1}^{\approx}(\mathcal{O}_1, \mathcal{O}_{\mathcal{O}_1, \mathcal{O}_2}^{\mathcal{M}})$ and $\text{diff}_{\Sigma_2}^{\approx}(\mathcal{O}_2, \mathcal{O}_{\mathcal{O}_1, \mathcal{O}_2}^{\mathcal{M}})$ must be empty for signatures $\Sigma_1 = \text{Sig}(\mathcal{O}_1)$ and $\Sigma_2 = \text{Sig}(\mathcal{O}_2)$, respectively.

In this paper we present two less restrictive variants of this principle: *subsumption* and *equivalence* conservativity violations. The *subsumption* variant of the conservativity principle requires $\mathcal{O}_{\mathcal{O}_1, \mathcal{O}_2}^{\mathcal{M}}$ not to introduce new subsumption relationships between concepts from one of the input ontologies, unless they were already involved in a subsumption relationship or they shared a common descendant. In addition, we also define violations between concepts that may have been already involved in a subsumption relationship (i.e., resulting in an equivalence between them), denoted as *equivalence conservativity principle violations*, or simply *equivalence violations*. Both variants are formally introduced in Definition 2.6.

Definition 2.6. Let $\mathcal{O}_i$ be one of the input ontologies and $\Sigma = \text{Sig}(\mathcal{O}_i) \setminus \{\perp, \top\}$ be its signature, let $\mathcal{M}$ be a *coherent* alignment between $\mathcal{O}_1$ and $\mathcal{O}_2$, $\mathcal{O}_{\mathcal{O}_1, \mathcal{O}_2}^{\mathcal{M}}$ be the integrated ontology,⁶ and let A, B be concepts in Σ . We define two sets of violations of $\mathcal{O}_{\mathcal{O}_1, \mathcal{O}_2}^{\mathcal{M}}$ w.r.t. $\mathcal{O}_i$:

- *subsumption violations*, denoted as $\text{subViol}(\mathcal{O}_i, \mathcal{O}_{\mathcal{O}_1, \mathcal{O}_2}^{\mathcal{M}})$, as the set of $A \sqsubseteq B$ axioms satisfying: (i) $A \sqsubseteq B \in \text{diff}_{\Sigma}^{\approx}(\mathcal{O}_i, \mathcal{O}_{\mathcal{O}_1, \mathcal{O}_2}^{\mathcal{M}})$, (ii) $\mathcal{O}_i \not\models B \sqsubseteq A$, and (iii) there is no C in Σ such that $\mathcal{O}_i \models C \sqsubseteq A$, and $\mathcal{O}_i \models C \sqsubseteq B$;
- *equivalence violations*, denoted as $\text{eqViol}(\mathcal{O}_i, \mathcal{O}_{\mathcal{O}_1, \mathcal{O}_2}^{\mathcal{M}})$, as the set of $A \equiv B$ axioms satisfying: (i) $\mathcal{O}_{\mathcal{O}_1, \mathcal{O}_2}^{\mathcal{M}} \models A \equiv B$, (ii) $A \sqsubseteq B \in \text{diff}_{\Sigma}^{\approx}(\mathcal{O}_i, \mathcal{O}_{\mathcal{O}_1, \mathcal{O}_2}^{\mathcal{M}})$ and/or $B \sqsubseteq A \in \text{diff}_{\Sigma}^{\approx}(\mathcal{O}_i, \mathcal{O}_{\mathcal{O}_1, \mathcal{O}_2}^{\mathcal{M}})$.

Thus, in our setting, an alignment $\mathcal{M}$ violates the conservativity principle if the sets $\text{subViol}(\mathcal{O}_i, \mathcal{O}_{\mathcal{O}_1, \mathcal{O}_2}^{\mathcal{M}})$ or $\text{eqViol}(\mathcal{O}_i, \mathcal{O}_{\mathcal{O}_1, \mathcal{O}_2}^{\mathcal{M}})$ are not empty. Example 2.1 provides an example of mappings leading to *subsumption* and *equivalence* conservativity violations.

Example 2.1. Consider the alignment $M = \{m_1, m_2, m_3, m_4\}$ provided in Figure 1. We distinguish the following cases depending on the semantic relation of the mappings:

- (i) *Both types of violation.* If both m_1 and m_2 are equivalence mappings (i.e., $\equiv$) then m_1 and m_2 lead to two *subsumption* violations and one *equivalence* violation since they will make *Joint₂* and *Set_of_Joints₂* to become equivalent, which did not hold any relationship in the original ontology.
- (ii) *Only subsumption violation.* In the case m_1 is an equivalence mapping (i.e., $\equiv$) but m_2 is a subsumption mapping (i.e., *Joint_Structure₁* $\sqsubseteq$ *Set_of_Joints₂*), m_1 and m_2 only lead to a *subsumption* violation (i.e., *Joint₂* $\sqsubseteq$ *Set_of_Joints₂*).
- (iii) *Only equivalence violation.* If both m_3 and m_4 are equivalence mappings (i.e., $\equiv$) then these mappings lead to the equivalence of *Anatomical_Structure₁* and *Musculoskeletal_System_Struct₁*, which represents an *equivalence* violation; however it does not represent a *subsumption* violation, according to condition (iii) in Definition 2.6, since *Joint_Structure₁* is a common descendant of the concepts involved in the novel subsumptions.

⁶ We assume that $\text{diff}_{\Sigma}^{\approx}(\mathcal{O}_i, \mathcal{O}_{\mathcal{O}_1, \mathcal{O}_2}^{\emptyset}) = \emptyset$.

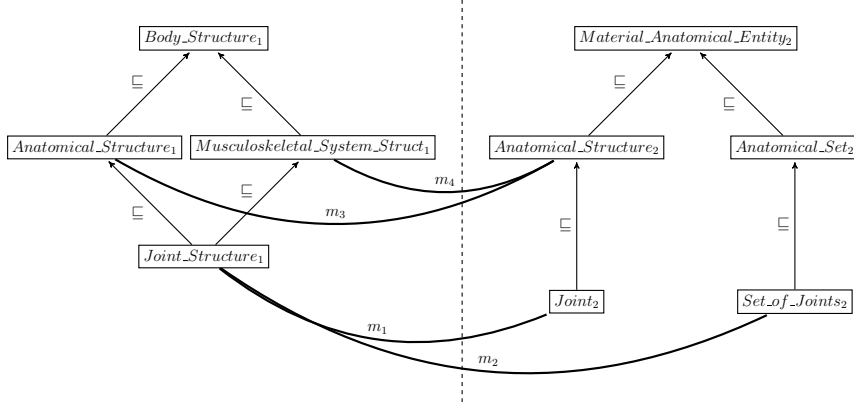


Fig. 1. Example of mappings leading to *subsumption* and *equivalence* conservativity violations. The ontology in the left-hand side represents a fragment of the *SNOMED CT* [69] ontology, while the one in the right-hand side represents a fragment of *FMA* [65] ontology.

- (iv) *No violations.* If m_3 is a equivalence mapping (i.e., $\equiv$) but m_4 is a subsumption mapping (i.e., $\sqsubseteq$), then m_3 and m_4 will only lead to the novel subsumption $Musculoskeletal_System_Struct_1 \sqsubseteq Anatomical_Structure_1$, which does not represent a *subsumption* violations for the same reason as described in the previous case.

The (potential) complexity of detecting and solving conservativity violations requires to exploit the intrinsic characteristics of the above two flavours, that will result in the development of different methods for their repair. For example, the notion of subsumption violations relies on the assumption of disjointness [67] and it can be reduced to a consistency principle problem (see Section 5). Equivalence violations, however, rely on the notion of (unsafe) cycle in the ontology graph and require a different treatment based on graph theory (Section 4).

Alignment Repair. An alignment $\mathcal{M}$ that violates the consistency and/or the conservativity principles can be fixed by removing correspondences from $\mathcal{M}$.⁷ This process is referred to as *mapping repair* (or repair for short), and is formally introduced in Definitions 2.7 and 2.8.

Definition 2.7. Let $\mathcal{M}$ be an incoherent alignment (i.e., violates the consistency principle) w.r.t. $\mathcal{O}_1$ and $\mathcal{O}_2$. A set of mappings $\mathcal{R} \subseteq \mathcal{M}$ is a *mapping repair for the consistency violations* for $\mathcal{M}$ w.r.t. $\mathcal{O}_1$ and $\mathcal{O}_2$ iff $\mathcal{M} \setminus \mathcal{R}$ is coherent w.r.t. $\mathcal{O}_1$ and $\mathcal{O}_2$, that is, $\mathcal{O}_{\mathcal{O}_1, \mathcal{O}_2}^{\mathcal{M} \setminus \mathcal{R}} \not\models A \sqsubseteq \perp$, for all $A \in Sig(\mathcal{O}_1) \cup Sig(\mathcal{O}_2)$.

The definition of violation of the conservativity principle, as in Definition 2.6, assumes that the alignment $\mathcal{M}$ does not violates the consistency principle w.r.t. the input ontologies $\mathcal{O}_1$ and $\mathcal{O}_2$. The main reason for requiring as input a coherent alignment is

⁷ Note that in this paper we only target the mappings in the repair process and we consider the input ontologies as immutable. Other approaches like Pesquita et al. [7] question the automatic generation of repairs and suggest to update the ontologies, when necessary, to avoid violations.

that unsatisfiable concepts would be subsumed by any other concept, and thus leading to a very large number of (misleading) violations of the conservativity principle.

Definition 2.8. Let $\mathcal{M}$ be a coherent alignment that violates the conservativity principle w.r.t. $\mathcal{O}_1$ and $\mathcal{O}_2$. A set of mappings $\mathcal{R} \subseteq \mathcal{M}$ is a *mapping repair for the conservativity violations* for $\mathcal{M}$ w.r.t. $\mathcal{O}_1$ and $\mathcal{O}_2$ iff for all $i \in \{1, 2\}$ $\text{subViol}(\mathcal{O}_i, \mathcal{O}_{\mathcal{O}_1, \mathcal{O}_2}^{\mathcal{M} \setminus \mathcal{R}})$ and $\text{eqViol}(\mathcal{O}_i, \mathcal{O}_{\mathcal{O}_1, \mathcal{O}_2}^{\mathcal{M} \setminus \mathcal{R}})$ are empty.

The repair of an alignment usually has several alternatives. Nevertheless, the objective is to remove as few mappings as possible. The notion of minimal repairs, originally introduced by Reiter [63] under the name of *diagnoses*, have been introduced to the field of ontology debugging in [68]. In the ontology debugging literature, they are typically referred to as *mapping diagnoses* [54]. A repair or diagnosis for an alignment can be computed by extracting the justifications for the unsatisfiable concepts (e.g., [29, 42, 67, 80]), and selecting a hitting set of mappings to be removed, following a minimality criteria (e.g., the number of removed mappings, or their combined confidence). However, justification-based techniques do not scale when the number of unsatisfiabilities is large (a typical scenario in mapping repair problems [34, 76, 77]). To address this scalability issue, mapping repair systems usually compute an *approximated repair* using incomplete reasoning techniques (e.g., [33, 54, 66]). An approximated repair $\mathcal{R}^\approx$ does not guarantee that $\mathcal{M} \setminus \mathcal{R}^\approx$ does not lead to violations of the consistency and conservativity principle, but it will (in general) significantly reduce the number of violations caused by the original set of mappings $\mathcal{M}$.

2.4. Graph Theory Notation

This section provides the necessary preliminaries about graph theory, given that graphs will be used to provide a convenient representation of ontologies for reasoning about the detection and repair of conservativity violations.

Graph Notions. A *directed graph* (digraph) G is a set of vertices V together with a relation A on V .⁸ A *weighted digraph* G is a digraph where each arc has a third component, called *weight* of the arc, assuming values in an appropriate structure (e.g., $\mathbb{N}, \mathbb{R}$, etc.). Given an arc $a = (u, v, c)$, we denote its weight as $w(a) = c$. The function w , when applied to a set of arcs, returns the sum of the weights of the single arcs. An *arc-labeled digraph* $G = (V, A, L)$ is a digraph with an additional component L called the set of labels. Given an arc $a = (u, v, l)$, l is called the label of a , where $l \in L$.

If not differently stated, in the remainder we always refer to *weighted digraphs*, and the placeholder “ $-$ ” might replace any of the components of an arc, representing an unspecified legal value that the considered component can assume.

Given a digraph $G = (V, A)$, a *subgraph* $G' = (V', A')$ of G is a digraph such that $V' \subseteq V$, and $A' \subseteq A$, and for all arc $a = (u, v, -) \in A'$, we have that $u, v \in V'$. A *subgraph* of G is said to be *induced* by V' , and denoted as $G|_{V'}$, if, for all pair of vertices $u, v \in V'$, $(u, v, -) \in A'$ iff $(u, v, -) \in A$.

Paths and Cycles. We now formally introduce the notion of path and cycles, that will be used to model subsumption and equivalence between concepts.

Given a digraph G , a (directed) *path* $\pi = [v_1, \dots, v_n]$ of G , with $n > 1$, is a

⁸ In our setting A is required to be antireflexive as we disallow self-arcs.

sequence of vertices where each pair of vertices v_i and v_{i+1} , with $i \in 1 \dots n - 1$, is connected by an arc $(v_i, v_{i+1}, -) \in A$. The *length* of such a path π is $n - 1$. Given two vertices $u, v \in V$, v is *reachable* from u iff a path π of length $m - 1$ exists such that, for some $m > 2$, $v_1 = u$ and $v_m = v$.

Given a digraph G and a directed path $\pi = [v_1, \dots, v_n]$ of G , we define π as a directed cyclic path (*cycle* in what follows) iff the first and last vertices coincide, i.e., $v_1 = v_n$. We say that a cycle κ is *broken* by the removal of any of its arcs (resp. vertices). We also say that a set of arcs (resp. vertices) breaks a set of cycles iff any of the cycles contains at least one element of the set. Note that we do not consider self-arcs, and therefore cycles with length equal to 1 are not allowed. Given a digraph G and two cycles κ_1 and κ_2 of G , we define κ_1 and κ_2 as *distinct cycles* if they are not a cyclic permutation one of the other.

Using the notion of cycle, we introduce the notion of *directed acyclic graph* (DAG).

Given an arc $a = (u, v, -)$, we refer to $\mathcal{V}(a) = \{u, v\}$ as the set containing its source and target vertices. $\mathcal{V}$ naturally extends to sets (resp. sequences) of arcs, as the union of its application on each element of the set (resp. sequence).

Given a path $\pi = [v_1 \dots v_n]$, we refer to $\mathcal{A}(\pi) = \{(v_i, v_{i+1}, -) \mid i \in [1 \dots n - 1]\}$ as the set containing all of the arcs of π . Similarly, given a set of vertices V' and a digraph $G = (V, A)$ such that $V' \subseteq V$, we refer to $\mathcal{A}(V') = \{(u, v, -) \in A \mid u, v \in V'\}$ as the set containing all the arcs of G between vertices of V' .

Graph Connectivity. Given that paths and cycles are used to encode subsumption and equivalence relations between concepts, graph connectivity represents a natural option for reasoning about transitivity of such relations. The notion of *strongly connected component* (SCC), for instance, splits the graph into equivalence concepts, thus identifying sets of equivalent concepts.

More formally, given a digraph $G = (V, A)$, a SCC of G is a maximal set of vertices $C \subseteq V$ such that for all $u, v \in C$, both u is reachable from v and viceversa. Notice that at least a cycle containing all the elements of the SCC exists, so cycle detection in a digraph G can be reduced to the identification of the SCCs of G . The set of SCCs of a digraph $G = (V, A)$ is denoted as $SCC(G)$. If $SCC(G) = \{V\}$ then G is a *strongly connected digraph*.

Tarjan's Algorithm [81] (Tarjan) finds the SCCs of a digraph $G = (V, A)$ using a single depth-first search with a time complexity in $\mathcal{O}(|V| + |A|)$.

Given a digraph $G = (V, A)$, the *Feedback Edge Set (FES)* problem aims at selecting a subset of A , called *feedback (edge) set*, with minimum cardinality, whose removal makes G acyclic. This problem is known to be *NP-hard* [18], as well as its weighted variant, *Weighted Feedback Edge Set (WFES)*. *WFES* differs from *FES* for what concerns the minimization objective for the feedback set, that is required to be minimal w.r.t. the sum of the weights of its elements. *FES* is also equivalent to the *Feedback Vertex Set (FVS)* problem [18], where a subset of V that makes G acyclic is sought.

3. Motivating Example

In this section, we show the problems caused by the violation of the conservativity principle when integrating ontologies via mappings in a real-world scenario. To this end, we consider, as motivating example, a use case based on the *Optique* project application domain [24, 44].⁹ *Optique* aims at facilitating scalable end-user access to big data

⁹ www.optique-project.eu

Ontology $\mathcal{O}_1$	Ontology $\mathcal{O}_2$
α_1 <i>WellBore</i> $\sqsubseteq \exists \text{belongsTo.Well}$	β_1 <i>Exploration.well</i> $\sqsubseteq \text{Well}$
α_2 <i>WellBore</i> $\sqsubseteq \exists \text{hasOperator.Operator}$	β_2 <i>Explor.borehole</i> $\sqsubseteq \text{Borehole}$
α_3 <i>WellBore</i> $\sqsubseteq \exists \text{locatedIn.Field}$	β_3 <i>Appraisal.exp.borehole</i> $\sqsubseteq \text{Explor.borehole}$
α_4 <i>AppraisalWellBore</i> $\sqsubseteq \text{WellBore}$	β_4 <i>Appraisal.well</i> $\sqsubseteq \text{Well}$
α_5 <i>ExplorationWellBore</i> $\sqsubseteq \text{WellBore}$	β_5 <i>Field</i> $\sqsubseteq \exists \text{hasFieldOperator.Field.operator}$
α_6 <i>Operator</i> $\sqsubseteq \text{Owner}$	β_6 <i>Field.operator</i> $\sqcap \text{Owner} \sqsubseteq \text{Field.owner}$
α_7 <i>Operator</i> $\sqsubseteq \text{Company}$	β_7 <i>Company</i> $\sqsubseteq \text{Field.operator}$
α_8 <i>Field</i> $\sqsubseteq \exists \text{hasOperator.Company}$	β_8 <i>Field.owner</i> $\sqsubseteq \text{Owner}$
α_9 <i>Field</i> $\sqsubseteq \exists \text{hasOwner.Owner}$	β_9 <i>Borehole</i> $\sqsubseteq \text{Continuant} \sqcup \text{Occurrent}$

Table 1. Simplified fragments of two ontologies in the oil and gas domain.

Mappings $\mathcal{M}$				
	e_1	e_2	n	ρ
m_1	$\mathcal{O}_1:\text{Well}$	$\mathcal{O}_2:\text{Well}$	0.9	$\equiv$
m_2	$\mathcal{O}_1:\text{WellBore}$	$\mathcal{O}_2:\text{Borehole}$	0.7	$\equiv$
m_3	$\mathcal{O}_1:\text{ExplorationWellBore}$	$\mathcal{O}_2:\text{Exploration.well}$	0.6	$\sqsubseteq$
m_4	$\mathcal{O}_1:\text{ExplorationWellBore}$	$\mathcal{O}_2:\text{Explor.borehole}$	0.8	$\equiv$
m_5	$\mathcal{O}_1:\text{AppraisalWellBore}$	$\mathcal{O}_2:\text{Appraisal.exp.borehole}$	0.7	$\equiv$
m_6	$\mathcal{O}_1:\text{Field}$	$\mathcal{O}_2:\text{Field}$	0.9	$\equiv$
m_7	$\mathcal{O}_1:\text{Operator}$	$\mathcal{O}_2:\text{Field.operator}$	0.7	$\sqsubseteq$
m_8	$\mathcal{O}_1:\text{Company}$	$\mathcal{O}_2:\text{Company}$	0.9	$\equiv$
m_9	$\mathcal{O}_1:\text{hasOperator}$	$\mathcal{O}_2:\text{hasFieldOperator}$	0.6	$\equiv$
m_{10}	$\mathcal{O}_1:\text{Owner}$	$\mathcal{O}_2:\text{Owner}$	0.9	$\equiv$

Table 2. Ontology mappings for the vocabulary in $\mathcal{O}_1$ and $\mathcal{O}_2$.

in the oil and gas industry. The project is focused around two demanding use cases provided by *Siemens* [45] and *Statoil* [44]. Optique advocates for an Ontology Based Data Access (OBDA) approach so that end-users formulate queries using the vocabulary of a domain ontology instead of composing queries directly against the database. Ontology-based queries (e.g., SPARQL queries) are then automatically rewritten to SQL and executed over the database (e.g., [64]). Ontology entities are linked to the database through *ontology-to-schema* mappings. Ontology-to-schema mappings are, however, out of the scope of the present study, and we therefore only focus on ontology-to-ontology mappings, or simply mappings (see Section 2.2).

Table 1 shows two simplified fragments of ontologies that are currently being used in the context of Optique. The ontology $\mathcal{O}_1$ has been directly bootstrapped from a relational database (e.g., [38]), and it is linked to the data through (direct) ontology-to-schema mappings. The ontology $\mathcal{O}_2$, instead, is a domain ontology based on the Norwegian Petroleum Directorate (NPD) FactPages¹⁰ [72], preferred by the Optique end-users to feed the Optique’s visual query formulation interface [79].¹¹

The integration via ontology alignment of $\mathcal{O}_1$ and $\mathcal{O}_2$ is required since the vocabulary in $\mathcal{O}_2$ is used to formulate queries (i.e., QF-Ontology), but only the vocabulary of $\mathcal{O}_1$ is connected to the database (i.e., DB-Ontology), as depicted in Figure 2. Consider the set of mappings $\mathcal{M}$ in Table 2 between $\mathcal{O}_1$ and $\mathcal{O}_2$ generated by an off-the-shelf ontology alignment system. Mappings are represented as 4-tuples (see Definition 2.1); for example the mapping m_2 suggests an equivalence relationship between the entities $\mathcal{O}_1:\text{WellBore}$ and $\mathcal{O}_2:\text{Borehole}$ with confidence 0.7.

¹⁰ <http://factpages.npd.no/factpages/>

¹¹ Optique uses OWL 2 QL ontologies for query rewriting, while the query formulation may be based on much richer OWL 2 ontologies. The axioms that fall outside the OWL 2 QL profile are either approximated or not considered for the rewriting.

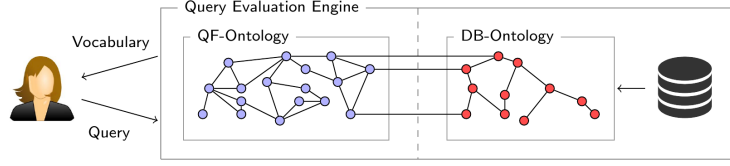


Fig. 2. Ontology Alignment in an OBDA Scenario

σ	Entailment:	follows from:	subViol?	eqViol?
σ_1	$\mathcal{O}_2:\text{Explor_borehole} \sqsubseteq \mathcal{O}_2:\text{Exploration_well}$	m_3, m_4	YES	NO
σ_2	$\mathcal{O}_1:\text{AppraisalWellBore} \sqsubseteq \mathcal{O}_1:\text{ExplorationWellBore}$	β_3, m_4, m_5	YES	NO
σ_3	$\mathcal{O}_2:\text{Field_operator} \sqsubseteq \mathcal{O}_2:\text{Field_owner}$	$\alpha_6, \beta_6, m_7, m_{10}$	YES	NO
σ_4	$\mathcal{O}_1:\text{Company} \equiv \mathcal{O}_1:\text{Operator}$	$\alpha_7, \beta_7, m_7, m_8$	NO	YES
σ_5	$\mathcal{O}_2:\text{Field_operator} \equiv \mathcal{O}_2:\text{Company}$			
σ_6	$\mathcal{O}_1:\text{Company} \sqsubseteq \mathcal{O}_1:\text{Owner}$	σ_4, α_6	YES	NO
σ_7	$\mathcal{O}_2:\text{Company} \sqsubseteq \mathcal{O}_2:\text{Field_owner}$	σ_3, σ_5	YES	NO

Table 3. Example of conservativity violations.

The integrated ontology $\mathcal{O}_{\mathcal{O}_1, \mathcal{O}_2}^{\mathcal{M}} = \mathcal{O}_1 \cup \mathcal{O}_2 \cup \mathcal{M}$, however, violates the conservativity principle (see Table 3). According to Definition 2.6, entailments σ_1 - σ_3 , σ_6 and σ_7 represent subsumption violations, while entailments σ_4 and σ_5 are equivalence violations. Note that σ_4 and σ_5 do not belong to the set of subsumption violations since $\mathcal{O}_1:\text{Company}$ and $\mathcal{O}_1:\text{Operator}$ (resp. $\mathcal{O}_2:\text{Field_operator}$ and $\mathcal{O}_2:\text{Company}$) are involved in a subsumption relationship in $\mathcal{O}_1$ (resp. $\mathcal{O}_2$).

In Figure 3 the portion of the graph representation of $\mathcal{O}_{\mathcal{O}_1, \mathcal{O}_2}^{\mathcal{M}}$ involved in conservativity violations is shown. Dashed arcs represent inferred axioms, while bold arcs are those involved in equivalence violations. Each non-inferred arc is labeled with its confidence value.

Example 3.1 provides an instance of query over the vocabulary of $\mathcal{O}_2$.

Example 3.1. Consider the following simple conjunctive query expressed in datalog notation, $CQ(x) \leftarrow \mathcal{O}_2:\text{Well}(x)$. The query asks for wells and has been formulated using the vocabulary of $\mathcal{O}_2$. The query is rewritten, according to the ontology axioms and mappings $\beta_1, \beta_4, m_1, m_3, m_4$ in $\mathcal{O}_{\mathcal{O}_1, \mathcal{O}_2}^{\mathcal{M}}$, into the following union of conjunctive

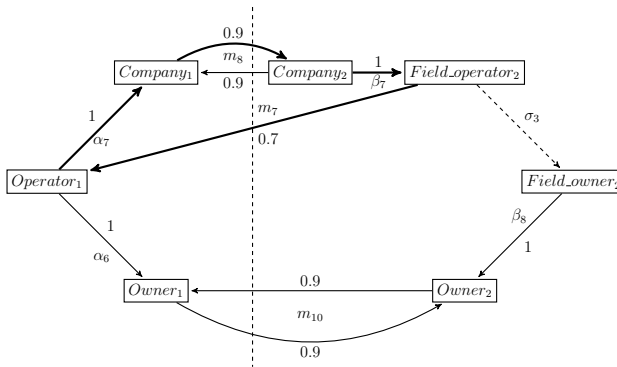


Fig. 3. Graph representation of the fragment of the aligned ontology involved in conservativity violations.

queries:

$$\begin{aligned} UCQ(x) \leftarrow & \mathcal{O}_2:Well(x) \cup \mathcal{O}_1:Well(x) \cup \mathcal{O}_2:Exploration_well(x) \cup \\ & \mathcal{O}_2:Appraisal_well(x) \cup \mathcal{O}_1:ExplorationWellBore(x) \cup \\ & \mathcal{O}_2:Explor_borehole(x) \end{aligned}$$

Since only the vocabulary of $\mathcal{O}_1$ is linked to the data, the union of conjunctive queries could be simplified as $UCQ(x) \leftarrow Well(x) \cup ExplorationWellBore(x)$, which will clearly lead to non desired results. The original query was only asking for wells, while the rewritten query will also return data about exploration wellbores.

Example 3.1 shows that the quality of the mappings in terms of conservativity violations may directly affect the quality of the query results in an OBDA context. Therefore, the detection and repair of these violations arise as an important quality assessment step in Optique. Conservativity violations, however, may represent *false positives*, that is, entailments that bring new and valid knowledge. In these cases, the detection and (suggested) repair of conservativity violations can be seen as an input of a subsequent manual revision where the bootstrapped and the domain ontologies may be updated with novel subsumption and/or equivalence axioms. Note that changes in the bootstrapped ontology will necessarily imply a revision of the ontology-to-schema mappings. Conservativity violations, however, must always be avoided in the case the bootstrapped and the domain ontologies are considered as immutable by Optique end-users.

4. Detecting and Repairing Equivalence Violations

This section introduces the relevant properties of our aligned ontology graph representation as well as the formal definition and analysis of the problem of computing a minimal diagnosis to repair equivalence violations. Finally we present our repair algorithm EqRepair which relies on logic programming.

As stated in Definitions 2.6 and 2.8, we expect as input an alignment $\mathcal{M}$ that is coherent w.r.t. the input ontologies $\mathcal{O}_1$ and $\mathcal{O}_2$.

4.1. Aligned Ontology and Graph Representation

Definition 4.1 formalizes a variant of the aligned ontology given in Definition 2.3. This variant enables a more efficient detection phase for equivalence violations. Differently from Definition 2.3, here we perform a classification step on the input ontologies before computing the aligned ontology, in order to obtain a graph representation that allows the detection of equivalence violations by simply analysing the graph, without further use of an OWL 2 reasoner.

Definition 4.1. Given two (input) ontologies, namely $\mathcal{O}_1$, $\mathcal{O}_2$, and an alignment $\mathcal{M}$ between them, consider an ontology $\mathcal{O}_{\mathcal{O}_1, \mathcal{O}_2}^{\mathcal{M}}$ such that $Sig(\mathcal{O}_{\mathcal{O}_1, \mathcal{O}_2}^{\mathcal{M}}) = Sig(\mathcal{O}_1) \cup Sig(\mathcal{O}_2)$, and $\mathcal{O}_{\mathcal{O}_1, \mathcal{O}_2}^{\mathcal{M}} = Cl(\mathcal{O}_1) \cup Cl(\mathcal{O}_2) \cup \mathcal{M}$, where Cl represents the classification of the given ontology as introduced in Section 2.1. We define $\mathcal{O}_{\mathcal{O}_1, \mathcal{O}_2}^{\mathcal{M}}$ as the *aligned ontology* w.r.t. $\mathcal{O}_1$, $\mathcal{O}_2$, and $\mathcal{M}$, or simply the aligned ontology, when no confusion arises.

Note that we do not compute the closure of $\mathcal{O}_1 \cup \mathcal{O}_2 \cup \mathcal{M}$ since this will prevent our repair algorithm to isolate the axioms (*i.e.*, mappings) causing an equivalence violation.

An example of aligned ontology is shown in Example 4.1.

$Cl(\mathcal{O}_1)$	$Cl(\mathcal{O}_2)$	$\mathcal{M}$
$Laptop \sqsubseteq PC$		
$Ultrabook \sqsubseteq Laptop \sqcap LightObject$	$Ultrabook \sqsubseteq Notebook$	$\langle Laptop_1, Notebook_2, \sqsupseteq, 0.7 \rangle$
$Ultrabook \sqsubseteq Laptop^*$	$Notebook \sqsubseteq Computer$	$\langle Ultrabook_1, Ultrabook_2, \equiv, 1 \rangle$
$Ultrabook \sqsubseteq LightObject^*$	$Ultrabook \sqsubseteq Computer^*$	$\langle PC_1, Computer_2, \sqsubseteq, 0.9 \rangle$
$Ultrabook \sqsubseteq PC^*$		

Table 4. Ontologies and alignment for Example 4.1. Inferred axioms are marked with an asterisk.

Example 4.1. Let $\mathcal{O}_1$ and $\mathcal{O}_2$ be two input ontologies and $\mathcal{M}$ be an alignment between them. Table 4 shows the classification of the input ontologies (inferred axioms are marked with an asterisk), and the mappings. The aligned ontology $\mathcal{O}_{\mathcal{O}_1, \mathcal{O}_2}^{\mathcal{M}}$, is therefore equal to $\{Laptop_1 \sqsubseteq PC_1, Ultrabook_1 \sqsubseteq Laptop_1 \sqcap LightObject_1, Ultrabook_1 \sqsubseteq Laptop_1, Ultrabook_1 \sqsubseteq LightObject_1, Ultrabook_1 \sqsubseteq PC_1, Ultrabook_2 \sqsubseteq Notebook_2, Notebook_2 \sqsubseteq Computer_2, Ultrabook_2 \sqsubseteq Computer_2, Notebook_2 \sqsubseteq Laptop_1, Ultrabook_1 \equiv Ultrabook_2, PC_1 \sqsubseteq Computer_2\}$. Note that subscripts are only used to emphasize the “provenance” of named concepts.

Note that our repair algorithm does not directly compute a diagnosis using the aligned ontology, it rather works on a graph representation¹² of this ontology, presented in Definition 4.2.

Definition 4.2. Given an ontology $\mathcal{O}$, its *graph representation*, denoted as $G(\mathcal{O}) = (V, A)$, is a digraph characterized by the following properties:

- each concept $C \in N_C(\mathcal{O})$ has an associated vertex $v_C \in V$,
- each axiom of the form $D \sqsubseteq E \in \mathcal{O}$ (resp. $D \sqsupseteq E$), with $D, E \in N_C(\mathcal{O})$ distinct concepts, has an associated arc (v_D, v_E, c) (resp. (v_E, v_D, c)) in A ,
- each axiom of the form $D \equiv E \in \mathcal{O}$, with $D, E \in N_C(\mathcal{O})$ distinct concepts, has an associated pair of arcs $\{(v_D, v_E, c), (v_E, v_D, c)\}$ in A .

The third component of each arc (denoted as c above) is always equal to 1 if $\mathcal{O}$ is not an aligned ontology, otherwise it is equal to 1 for the axioms of the input ontologies, while it is equal to the confidence of the considered mapping for the corresponding arcs.

When no confusion arises we interchangeably refer to an ontological concept and its associated vertex, and to an ontology and its associated graph representation. Note that subsumption/equivalence axioms between a concept and itself are not allowed in the graph representation, because the set of arcs is defined using an antireflexive relation. Similarly, we will interchangeably refer to axioms/mappings and their corresponding elements in the graph representation.

An example of graph representation, associated with the aligned ontology of Example 4.1, is shown in Figure 4.

4.2. Paths and Cycles

Given the semantics of the arcs in the graph representation, a path from a vertex u to a vertex v implies that $\mathcal{O}$ entails that the concept associated with u is subsumed by

¹² Despite several proposals for graph formalisms for representing DL ontologies exists in literature (e.g., [58]), we provided a simplified variant specifically tailored to capture equivalence violations.

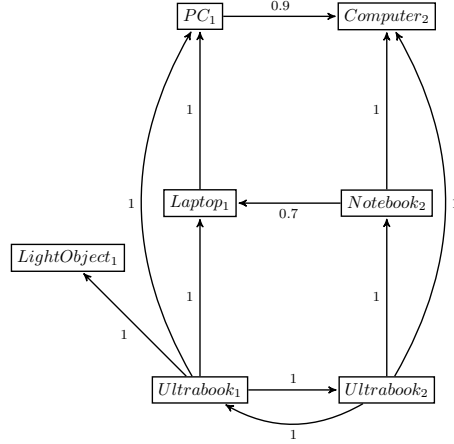


Fig. 4. Graph representation for the aligned ontology of Example 4.1.

the concept associated with v , as expressed by Proposition 4.1. Analogously, Proposition 4.2 states that a cycle in the graph representation implies that $\mathcal{O}$ entails that all the concepts associated with the vertices of the cycle are equivalent. Note that the other direction of the following propositions holds only if the graph representation is built on top of an ontology closed w.r.t. classification, but this is not the case for our technique. Therefore, we do not require that the graph representation of the aligned ontology reflects all its subsumption entailments (equivalently, it suffices that testing subsumption on the graph representation is sound, even if not complete).

Proposition 4.1. Let $\mathcal{O}$ be an ontology and let $G(\mathcal{O}) = (V, A)$ be its graph representation. If a path $\pi = [v_{A_1}, \dots, v_{A_n}]$ with $n > 1$ exists in $G(\mathcal{O})$, then $\mathcal{O} \models A_1 \sqsubseteq A_n$, with A_1, A_n distinct concepts.

Proof. The proof directly follows from the semantics of the arcs in the graph representation and by transitivity of the subsumption relation. $\square$

Proposition 4.2. Let $\mathcal{O}$ be an ontology and let $G(\mathcal{O}) = (V, A)$ be its graph representation. If a cycle $\kappa = [v_{A_1}, \dots, v_{A_n}, v_{A_1}]$ with $n \geq 1$ exists in $G(\mathcal{O})$, then $\mathcal{O} \models A_i \equiv A_j$, with $i, j \in [1 \dots n]$ and A_i, A_j distinct concepts.

Proof. Given a pair of vertices v_{A_i}, v_{A_j} , we define two paths $\pi_{ij} = [v_{A_i}, \dots, v_{A_j}]$ and $\pi_{ji} = [v_{A_j}, \dots, v_{A_i}]$ such that vertices v_{A_p}, v_{A_q} are consecutive in π_{ij}, π_{ji} , only if the same holds in κ , with $i, j, p, q \in [1 \dots n]$. Given that all the π_{ij} and π_{ji} are paths, we can apply Proposition 4.1. $\square$

Safe Cycles. In Definition 4.3 safe cycles are introduced, that is, cycles of the aligned ontology that do not violate the conservativity principle. This notion is key to discriminate between cycles representing or not an equivalence violation, and it is also at the basis of the encoding of equivalence violations in terms of our graph representation. Intuitively, a cycle is safe if the projection on each of the two input ontologies is either a single vertex, or a cycle traversing all the vertices of such projection exists. A special case of safe cycles is represented by cycles already existing in the input ontologies. The different kinds of safe cycles stem from the combination of the following conditions on the two projections of the cycle over the input ontologies: a projection is empty (condi-

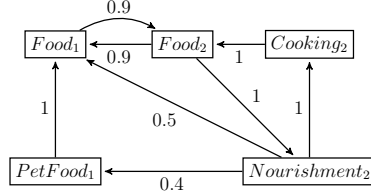


Fig. 5. A graph representation including both safe and unsafe cycles.

tion (i)), a projection consists of a single vertex (condition (ii)), a cycle traversing all the vertices of a projection exists (condition (iii)). For instance, a (safe) cycle is local to the input ontology $\mathcal{O}_i$ if its projection on $\mathcal{O}_i$ satisfies condition (iii), while the other projection is empty (condition (i) holds). This implies an equivalence entailment (involving all the concepts of the cycle) already in the input ontology $\mathcal{O}_i$, that clearly prevents the existence of an equivalence violation among such concepts.

Definition 4.3. Let $\kappa = [u_1, \dots, u_n, u_1]$ be one of the cycles of a graph representation $G(\mathcal{O}_{\mathcal{O}_1, \mathcal{O}_2}^{\mathcal{M}})$, with $n > 1$. We define κ as a *safe cycle* iff for some $i \in \{1, 2\}$ we have that: either (i) $|\mathcal{V}(\kappa) \cap V_{\mathcal{O}_i}| = 0$, or (ii) $|\mathcal{V}(\kappa) \cap V_{\mathcal{O}_i}| = 1$, or (iii) a cycle κ'_i in $G(\mathcal{O}_{\mathcal{O}_1, \mathcal{O}_2}^{\mathcal{M}})|_{V_{\mathcal{O}_i}}$ exists, such that $\mathcal{V}(\kappa) \cap V_{\mathcal{O}_i} \subseteq \mathcal{V}(\kappa'_i)$. We further differentiate between *local* ((i) and (iii) hold), *trivial* (only (ii) holds), *partially trivial* ((ii) and (iii) hold) and *nontrivial safe cycles* (only (iii) holds). We generically refer to safe cycles, except local safe cycles, as *global safe cycles*. A cycle that is not safe is defined as an *unsafe cycle*.

In Example 4.2 an example is provided for each different kind of cycles.

Example 4.2. Consider the graph representation of Figure 5. We first present the (different kind of) safe cycles. $[Food_2, Nourishment_2, Cooking_2, Food_2]$ is a local safe cycle, $[Food_1, Food_2, Nourishment_2, Food_1]$ is partially trivial, $[Food_1, Food_2, Food_1]$ is trivial. Instead, $[Food_2, Nourishment_2, PetFood_1, Food_1, Food_2]$ is an unsafe cycle, but adding the axiom $Food_1 \sqsubseteq PetFood_1$ to $\mathcal{O}_1$ would turn it into a nontrivial safe cycle.

Starting from the results of Proposition 4.2, we can characterize a restricted version of the conservativity principle using graph-theoretical concepts only, applied on the graph representation, without the need to refer to the aligned ontology. Despite the aforementioned approximations, our repair technique for equivalence violations is effective in practice, as experimentally verified in Section 7.2.

Theorem 4.1. Let $\mathcal{O}_1, \mathcal{O}_2$ be two (input) ontologies, and let $\mathcal{M}$ be an alignment between them. Let also $G(\mathcal{O}_{\mathcal{O}_1, \mathcal{O}_2}^{\mathcal{M}}) = (V, A)$, $G(\mathcal{O}_1) = (V_{\mathcal{O}_1}, A_{\mathcal{O}_1})$ and $G(\mathcal{O}_2) = (V_{\mathcal{O}_2}, A_{\mathcal{O}_2})$ be the graph representations associated with $\mathcal{O}_{\mathcal{O}_1, \mathcal{O}_2}^{\mathcal{M}}$, $\mathcal{O}_1$ and $\mathcal{O}_2$, respectively. If there exists an unsafe cycle κ^u in $G(\mathcal{O}_{\mathcal{O}_1, \mathcal{O}_2}^{\mathcal{M}})$ and the vertices associated with concepts C, D (namely V_C, V_D) belong to κ^u , then a violation of the conservativity principle w.r.t. equivalence, involving concepts C and D , exists in $\mathcal{O}_{\mathcal{O}_1, \mathcal{O}_2}^{\mathcal{M}}$.

Proof. For all the $v_C, v_D \in c$, if we have an (unsafe) cycle κ^u in $G(\mathcal{O}_{\mathcal{O}_1, \mathcal{O}_2}^{\mathcal{M}})$, then $\mathcal{O}_{\mathcal{O}_1, \mathcal{O}_2}^{\mathcal{M}} \models C \equiv D$ holds. For all $i \in \{1, 2\}$, if an unsafe cycle κ^u exists, then no cycle κ' exists in $G(\mathcal{O}_i)$ such that $\mathcal{V}(\kappa^u) \cap V_{\mathcal{O}_i} \subseteq \mathcal{V}(\kappa')$. This requires that at least one vertex $v_E \in \kappa^u$ exists such that it does not belong to κ' ; which corresponds to having that, for all $i \in \{1, 2\}$ and for some $E, F \in N_C(\mathcal{O}_i)$, $\mathcal{O}_i \not\models E \equiv F$ and $\mathcal{O}_{\mathcal{O}_1, \mathcal{O}_2}^{\mathcal{M}} \models E \equiv$

F . This represents an equivalence violation, thus proving the soundness of the detected violations using the graph representation. $\square$

Proposition 4.2 relates cycles in the graph representation to equivalence axioms in the aligned ontology. We have only soundness guarantee, given that without classifying the aligned ontology equivalences that are not reflected in the graph representation may exist.

Theorem 4.1 guarantees that a method detecting and correcting *all* the unsafe cycles on the graph representation is sound w.r.t. violations of the conservativity principle w.r.t. equivalence in the aligned ontology. This is proved by verifying that detected unsafe cycles effectively encode equivalence violations.

We remark also that, once the graph representation is built, we do not need the support of a standard reasoner. This feature is at the basis of the scalability of our approach.

4.3. Diagnoses for Equivalence Conservativity Violations

The goal of our repair algorithm is, starting from an aligned ontology, to detect violations to the conservativity principle w.r.t. equivalence and to compute a diagnosis. In order to obey to the well-known *principle of minimal change*, the diagnosis is required to have a minimal weight, w.r.t. a metric capturing the amount of lost information. In our context, a standard minimality criterion for defining diagnosis weight (the quantity to minimize) is the sum of the weights of the arcs associated with the (removed) mappings. In Definition 4.4 we formalize a diagnosis as the set of arcs of the graph representation of an aligned ontology that, once removed, breaks all the unsafe cycles.

Definition 4.4. Let $\mathcal{M}$ be an alignment such that $G(\mathcal{O}^{\mathcal{M}})$ has unsafe cycles $\{\kappa_1^u, \dots, \kappa_n^u\}$. Let also $\Delta \subseteq \mathcal{M}$ be an alignment. Δ is a *diagnosis* for $\mathcal{M}$ iff for all $i \in [1 \dots n]$, $\Delta \cap \mathcal{A}(\kappa_i^u) \neq \emptyset$, that is, the graph representation $G(\mathcal{O}^{\mathcal{M} \setminus \Delta})$ has no unsafe cycles.

Alternatively, we refer to a diagnosis for $\mathcal{M}$ as the diagnosis for the aligned ontology $\mathcal{O}^{\mathcal{M}}$ or its graph representation $G(\mathcal{O}^{\mathcal{M}})$. In Definition 4.5 we formalize diagnosis minimality (required by the principle of minimal change), where given a diagnosis $\Delta = \{a_0, \dots, a_n\}$, its weight, denoted as $w(\Delta)$, is equal to $\sum_{i=1}^n w(a_i)$.

Definition 4.5. Let Δ be a diagnosis for an alignment $\mathcal{M}$. Δ is a *minimal diagnosis* for $\mathcal{M}$ iff there is no diagnosis Δ' for $\mathcal{M}$ such that $w(\Delta') < w(\Delta)$.

Before introducing the notion of problematic SCC in Definition 4.6, key for the definition of an efficient detection method for unsafe cycles, we need to define the *projection of a SCC* $S \in SCC(G(\mathcal{O}_{\mathcal{O}_1, \mathcal{O}_2}^{\mathcal{M}}))$ on an input ontology $\mathcal{O}_i$, denoted as $\Pi_{\mathcal{O}_i}(S)$, as $S \cap V_{\mathcal{O}_i}$, where $G_{\mathcal{O}_i} = (V_{\mathcal{O}_i}, A_{\mathcal{O}_i})$ is the graph representation of $\mathcal{O}_i$. We denote with $SCC_{local}(G(\mathcal{O}_{\mathcal{O}_1, \mathcal{O}_2}^{\mathcal{M}}))$ the set of SCCs of $G(\mathcal{O}_{\mathcal{O}_1, \mathcal{O}_2}^{\mathcal{M}})$.

Definition 4.6. A SCC $S \in SCC(G(\mathcal{O}_{\mathcal{O}_1, \mathcal{O}_2}^{\mathcal{M}}))$ is *problematic* iff $\Pi_{\mathcal{O}_1}(S) \notin SCC_{local}(G(\mathcal{O}_{\mathcal{O}_1, \mathcal{O}_2}^{\mathcal{M}})) \vee \Pi_{\mathcal{O}_2}(S) \notin SCC_{local}(G(\mathcal{O}_{\mathcal{O}_1, \mathcal{O}_2}^{\mathcal{M}}))$. The set of *problematic SCCs* of $G(\mathcal{O}_{\mathcal{O}_1, \mathcal{O}_2}^{\mathcal{M}})$ is denoted as $pSCC(G(\mathcal{O}_{\mathcal{O}_1, \mathcal{O}_2}^{\mathcal{M}}))$.

Example 4.3. Consider the graph represented in Figure 6. It is a subgraph of the aligned ontology $\mathcal{O}^{UMLS}$ representing the *UMLS* 2012 alignment, namely *UMLS*, between ontologies $\mathcal{O}_1, \mathcal{O}_2$, corresponding to the *FMA* [65] and *NCI* [26]. The graph represents a problematic SCC because its projection on $\mathcal{O}_2$ does not belong to the local SCCs of $\mathcal{O}_2$.

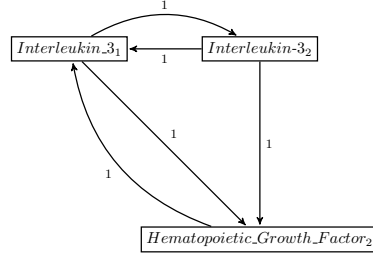


Fig. 6. Example of cycle defined by the *UMLS* 2012 alignment between *FMA* and *NCI* ontologies. The equivalence of the three vertices is entailed by the aligned ontology.

since no cycle between its two vertices exists. It cannot traverse another SCC, because otherwise the two SCCs would be a single one.

Example 4.4. Consider two input ontologies $\mathcal{O}_1 = \{MolarTooth_1 \sqsubseteq Tooth_1\}$ and $\mathcal{O}_2 = \{Tooth_{19_2} \sqsubseteq MolarTooth_2\}$ and their obvious associated signatures. Consider also the alignment $\{\langle MolarTooth_1, MolarTooth_2, \equiv, 1 \rangle, \langle Tooth_1, Tooth_{19_2}, \equiv, 0.7 \rangle\}$. The resulting digraph contains a SCC including all the four vertices, but none of the two projections on the input ontologies are SCCs. As a consequence, all these vertices are equivalent in the aligned ontology. From the domain knowledge, however, we know that $Tooth_1 \sqsubseteq MolarTooth_1$ and $Tooth_{19_2} \sqsubseteq MolarTooth_2$ do not hold.

The set of *problematic mappings* is the set of mappings between vertices of a problematic SCC, the only arcs that can appear in diagnoses.

The notion of problematic SCCs enables an alternative diagnosis definition. Using the notion of problematic SCC, Definition 4.7 and Definition 4.8 propose an alternative definition for diagnosis, w.r.t. that given in Definition 4.4.

Definition 4.7. Let Δ be an alignment such that $\Delta \subseteq \mathcal{M}$, and let $G(\mathcal{O}^{\mathcal{M} \setminus \Delta}) = (V, A')$ be a graph representation such that $A' = A \setminus \Delta$. Δ is a *diagnosis* for $\mathcal{M}$ w.r.t. $G(\mathcal{O}^{\mathcal{M}})$ iff $pSCC(G(\mathcal{O}^{\mathcal{M} \setminus \Delta})) = \emptyset$.

Definition 4.8. Let Δ be a diagnosis for $\mathcal{M}$, let $S \in pSCC(G(\mathcal{O}^{\mathcal{M}}))$ be a problematic SCC, and let Δ' be a nonempty subset of Δ . Δ' is a *diagnosis for a SCC* S iff for all S_i in $SCC(G(\mathcal{O}^{\mathcal{M} \setminus \Delta'}))$ such that $S_i \subseteq S$, we have that $S_i \notin pSCC(G(\mathcal{O}^{\mathcal{M} \setminus \Delta'}))$.

Clearly, the reformulation does not affect any of the properties of diagnosis illustrated in this section. Example 4.5 shows a minimal diagnosis for a problematic SCC.

Example 4.5. Consider the scenario of Example 4.3 and the subgraph of Figure 6, representing a problematic SCC. A minimal diagnosis for the problematic SCC is $\{(Interleukin_3_1, Hematopoietic_Growth_Factor_2, 1)\}$.

4.4. Problem Statement for Diagnosis Computation

This section introduces the *MAP-WFES* problems as an extension of the well-known *WFES* problem, which removes unsafe cycles only for solving the equivalence violations, while preserving local cycles. The computational complexity of the *MAP-WFES* problem is analyzed in Appendix A.1.

MAP-WFES Definition. Let $G = (V, A)$ be a digraph such that $V = V_1 \cup V_2$, with $V_1 \cap V_2 = \emptyset$. We denote $\mathcal{M}$, the set of *mappings*, the subset of A whose arcs are of the form $(u, v, -)$, such that for all $i, j \in \{1, 2\}$ and $i \neq j$ we have that $u \in V_i, v \in V_j$.

MAP-WFES aims at computing, given as input the digraph G , a minimum weight feedback edge set $\Delta \subseteq \mathcal{M}$ such that no unsafe cycles exist in $G' = (V, A \setminus \Delta)$.

Proposition 4.3 relates diagnosis computation to the *MAP-WFES* problem introduced in this section. Specifically, it states that computing a diagnosis for a graph representation reduces to solving *MAP-WFES* on it.

Proposition 4.3. Computing a minimal diagnosis for an aligned ontology $\mathcal{O}_{\mathcal{O}_1, \mathcal{O}_2}^{\mathcal{M}}$, that is, solving all the violations to the conservativity principle w.r.t. equivalence, can be reduced to solving an instance of *MAP-WFES*.

Proof. From Theorem 4.1 we know that all the violations of the conservativity principle w.r.t. equivalence in the aligned ontology result in unsafe cycles in its graph representation. Therefore breaking all unsafe cycles is equivalent to computing a diagnosis. From the definition of *MAP-WFES*, we have that no unsafe cycles can exist in the graph resulting from the application of the computed diagnosis. Given that, with the exception of computing a minimal diagnosis, *MAP-WFES* has no other constraints, it computes such a diagnosis, and this concludes the proof. $\square$

Since *MAP-WFES* is a *NP*-hard problem (as shown in Appendix A.1), and given the average size of an aligned ontology, computing a diagnosis would be, in practice, intractable. For this reason, our approach decomposes the problem into independent subproblems (*i.e.*, computing a local diagnosis for each problematic SCC), following a *divide et impera* strategy. A (minimal) global diagnosis is then computed from the (minimal) local diagnoses of the single problematic SCCs. The optimality and correctness of this decomposition are investigated in Appendix A.2.

4.5. Solving Equivalence Violations using Logic Programming

This section introduces the method for computing minimal diagnoses using logic programming, and a detailed description of the EqRepair algorithm.

Minimal Diagnosis Using ASP. We now describe how a minimal diagnosis can be computed using Answer Set Programming (*ASP*) programs.^{13 14} The facts input to the *ASP* problems are of the following kinds: (i) Vertices are represented using a binary predicate $vtx(X, O)$, where X is a string representing the vertex label and $O \in \{1, 2\}$ encodes the index of the input ontology the represented concept belongs to, (ii) arcs are represented using a quaternary predicate $edge(X, Y, C, M)$, where X, Y are vertices (that is, for some $O, P \in \{1, 2\}$, $vtx(X, O)$ and $vtx(Y, P)$ hold), C is an integer encoding the arc weight in a range $[0 \dots 100]$, while M is a Boolean flag for differentiating arcs that correspond to axioms ($M = 0$) from those corresponding to mappings ($M = 1$).

When we refer to the execution of an *ASP* program on a graph representation (resp.

¹³ We use the syntax of *Lparse 1.0*, a parser for logic programs used as a front-end by different logic programming solvers, more details at <http://www.tcs.hut.fi/Software/smodels/>

¹⁴ Despite alternative frameworks could have been employed (*e.g.*, Constraint Logic Programming), we have adopted *ASP* as it is known to be well-suited for graph-related problems, and to produce compact and easy to understand solutions [13].

Listing 1: *ASP* facts encoding the problematic SCC of Figure 3.

```

vtx(Company1,1). vtx(Operator1,1). vtx(Company2,2). vtx(Field_operator2,2).
edge(Company1,Company2,90,1). edge(Company2,Company1,90,1).
edge(Company2,Field_operator2,100,0). edge(Field_operator2,Operator1,70,1).

```

Listing 2: *ASP* program computing a minimal diagnosis for the *MAP-WFES* problem.

```

r0: #domain vtx(X,O). #domain vtx(Y,P). #domain vtx(Z,Q).
r1: reaches(X,Y) :- edge(X,Y,C,M), not_removed(edge(X,Y,C,M)), X!=Y.
r2: reaches(X,Z) :- reaches(X,Y), edge(Y,Z,C,M), not_removed(edge(Y,Z,C,M)),
    X!=Y, Y!=Z, X!=Z.
r3: reachesSafe(X,Y) :- edge(X,Y,C,0), O=P, X!=Y.
r4: reachesSafe(X,Z) :- reachesSafe(X,Y), edge(Y,Z,C,0), O=P, X!=Y, Y!=Z, X!=Z.
r5: not_removed(edge(X,Y,C,M)) | removed(edge(X,Y,C,M)) :- edge(X,Y,C,M), X!=Y.
r6: not_removed(edge(X,Y,C,0)) :- edge(X,Y,C,0), X!=Y, O=P.
r7: unsafeCycle(Y) :- not reachesSafe(Y,X), reaches(Y,X), reaches(X,Y), O=P, X!=Y.
r8: :- unsafeCycle(Y).
r9: #minimize [ removed(edge(X,Y,C,1)) = C ].
r10: #hide. #show removed/1.

```

an aligned ontology), we always implicitly refer to its execution on a set of facts encoding the graph representation (resp. aligned ontology).

Example 4.6. In the following, we provide an example of encoding of a graph into a set of facts that can be used in conjunction with the *ASP* programs for solving the *MAP-WFES* problems. To this end, we rely on our motivating example presented in Section 3. The encoding of the problematic SCC of Figure 3 corresponds to the set of facts presented in Listing 1. The execution of the *ASP* program of Listing 2, in conjunction with the aforementioned facts, is equal to the mapping m_7 from Table 2 (*i.e.*, the candidate mapping(s) to be removed to solve the violation(s)).

A minimal diagnosis can be computed using the *ASP* program shown in Listing 2. Rule r0 “types” variables X , Y or Z as vertex id (into the *vtx* predicate they appear). Rule r1 states that if an unremoved arc $(X, Y, -)$ exists, then vertex X *reaches* vertex Y , rule r2 states that *reaches* is a transitive predicate. Similarly, rule r3 that if an arc $(X, Y, -)$ exists, that is neither a mapping nor removed, then vertex X *reachesSafe* vertex Y , and rule r4 makes *reachesSafe* a transitive predicate.

In order to obtain a valid solution, each arc has to be either removed or not removed (rule r5), and that only mappings can be removed (rule r6). Additionally, rule r8 forbids the existence of unsafe cycles, identified by rule r7 (if a vertex unsafely reaches itself, then we have at least an unsafe cycle). Among the valid solutions, rule r9 minimizes diagnosis’ weight, and the output model is restricted to *removed* predicate only by means of rule r10, that corresponds to the computed diagnosis.

EqRepair Algorithm. The EqRepair algorithm¹⁵ (Algorithm 1) takes as input two ontologies $\mathcal{O}_1$ and $\mathcal{O}_2$, and a (coherent) alignment $\mathcal{M}$ between them. The graph representation G of the aligned ontology w.r.t. $\mathcal{O}_1$, $\mathcal{O}_2$ and $\mathcal{M}$, is built by means of createDigraph function (line 1 of Algorithm 1). The function builds a digraph G representing the aligned ontology associated with the input ontologies $\mathcal{O}_1$, $\mathcal{O}_2$ and alignment $\mathcal{M}$. In

¹⁵ In [73] this algorithm is referred to as CycleBreaker.

Algorithm 1 EqRepair Algorithm for *equivalence* conservativity violations**Input:** $\mathcal{O}_1, \mathcal{O}_2$: input ontologies; $\mathcal{M}$: input (coherent) mappings**Output:** Δ : diagnosis

```

1:  $G \leftarrow \text{createDigraph}(\mathcal{O}_1, \mathcal{O}_2, \mathcal{M})$ 
2:  $SCCS_i \leftarrow \text{tarjan}(G, \mathcal{O}_i)$ 
3:  $globalSCCs \leftarrow \text{tarjan}(G)$ 
4:  $mappings \leftarrow \emptyset$  ▷ Map of SCCs' mappings
5:  $\Delta \leftarrow \emptyset$  ▷ Diagnosis for  $\mathcal{M}$ 
6: for each  $S$  in  $globalSCCs$  do
7:   if  $\neg \bigwedge_{i=1}^2 \Pi_{\mathcal{O}_i}(S) \in SCCS_i$  then
8:      $mappings(S) \leftarrow \text{extractMappings}(S)$ 
9:      $\Delta \leftarrow \Delta \cup \text{solver}(S)$ 
10:   end if
11: end for
12: return  $\Delta$ 

```

accordance with Definition 4.2, the vertices of this graph are the named concepts of the two ontologies, and its arcs are the axioms/mappings involving them.¹⁶

After the creation of the digraph, the SCCs of the input ontologies (line 2) and that of the aligned ontology (line 3) are computed, by means of the *Tarjan's* algorithm (introduced in Section 2.4). The algorithm then detects which SCCs of the aligned ontology are problematic. By Definition 4.6, it suffices to test if at least one of the two projections on the input ontologies of each considered SCC is not a local SCC (line 7).

The minimal diagnosis for the current SCC, namely S , is computed by the solver function, that runs an *ASP* solver over a program representing an instance of the *MAP-WFES* problem associated with S (line 9 of Algorithm 1). The diagnosis for S is then obtained, using a translation, from the solution of the aforementioned problem. Theorem 4.1 and Proposition 4.3 detail, respectively, the correctness of the detection and the solution computation, as described above. The global diagnosis results from the union of the local diagnoses, that is, the diagnoses of the single SCCs (the correctness proof is given in Proposition A.5, Appendix A.2).

5. Detecting and Repairing Subsumption Violations

Our technique for coping with subsumption violations is based on the reduction of the conservativity principle repair problem to a consistency principle repair problem (*i.e.*, a mapping incoherence repair problem) through the assumption of disjointness [67]. Currently, our method reuses and adapts the structural indexing and reasoning techniques implemented in LogMap (an ontology matching and mapping repair system [33, 37, 39]). However, alternative mapping repair systems could be used, such as ALCOMO [54] or AML [19]. Note that, to the best of our knowledge, these mapping repair systems have only focused on solving violations of the consistency principle.

Algorithm 2 shows *SubRepair* algorithm, the proposed method for detecting and correcting subsumption violations. SubRepair algorithm expects as input an alignment $\mathcal{M}$ that is coherent w.r.t. the input ontologies $\mathcal{O}_1$ and $\mathcal{O}_2$, according to Definitions 2.6 and 2.8. SubRepair algorithm outputs the number of added disjointness during the process *disj* and an (approximate) repair $\mathcal{R}^\approx$. The following paragraphs describe the techniques used at each step.

¹⁶ For sake of space the algorithm, which directly follows from Definition 4.2, is omitted. The interested reader can find it in [73], Algorithm 12.

Algorithm 2 SubRepair algorithm for *subsumption* conservativity violations**Input:** $\mathcal{O}_1, \mathcal{O}_2$: input ontologies; $\mathcal{M}$: input (coherent) mappings**Output:** $\mathcal{R}^\approx$: approximate repair; $disj$: number of disjointness rules

```

1:  $\langle \mathcal{P}_1, \mathcal{P}_2 \rangle \leftarrow \text{propositionalEncoding}(\mathcal{O}_1, \mathcal{O}_2)$ 
2:  $SI_1 \leftarrow \text{structuralIndex}(\mathcal{O}_1)$ 
3:  $SI_2 \leftarrow \text{structuralIndex}(\mathcal{O}_2)$ 
4:  $\mathcal{O}_{\mathcal{O}_1, \mathcal{O}_2}^\mathcal{M} \leftarrow \mathcal{O}_1 \cup \mathcal{O}_2 \cup \mathcal{M}$  ▷ The aligned ontology is computed
5:  $SI_\mathcal{U} \leftarrow \text{structuralIndex}(\mathcal{O}_1 \cup \mathcal{O}_2 \cup \mathcal{M})$ 
6:  $\langle \mathcal{P}_1^d, disj_1 \rangle \leftarrow \text{disjointAxiomsExtension}(\mathcal{P}_1, SI_1, SI_\mathcal{U}, \mathcal{O}_{\mathcal{O}_1, \mathcal{O}_2}^\mathcal{M})$  ▷ See Algorithm 3
7:  $\langle \mathcal{P}_2^d, disj_2 \rangle \leftarrow \text{disjointAxiomsExtension}(\mathcal{P}_2, SI_2, SI_\mathcal{U}, \mathcal{O}_{\mathcal{O}_1, \mathcal{O}_2}^\mathcal{M})$ 
8:  $\langle \mathcal{M}', \mathcal{R}^\approx \rangle \leftarrow \text{mappingRepair}(\mathcal{P}_1^d, \mathcal{P}_2^d, \mathcal{M})$  ▷ See Algorithm 2 in [39]
9:  $disj \leftarrow disj_1 + disj_2$ 
10: return  $\langle \mathcal{R}^\approx, disj \rangle$ 

```

Propositional Horn Encoding. The ontologies $\mathcal{O}_1$ and $\mathcal{O}_2$ are encoded as the Horn propositional formulas, $\mathcal{P}_1$ and $\mathcal{P}_2$ (line 1 in Algorithm 2). For example, the concept hierarchy provided by an OWL 2 reasoner (e.g., [25, 43]) are encoded as $A \rightarrow B$ rules, while the explicit disjointness relationships between concepts are represented as $A_i \wedge A_j \rightarrow \perp$. Note that the input mappings $\mathcal{M}$ can already be seen as propositional implications. This encoding is key to the mapping repair process.

Example 5.1. Consider the ontologies and mappings in Tables 1 and 2. The axiom β_6 is encoded as $Field_operator \wedge Owner \rightarrow Field_owner$, while the mapping m_2 is translated into rules $\mathcal{O}_1:WellBore \rightarrow \mathcal{O}_2:Borehole$, and $\mathcal{O}_2:Borehole \rightarrow \mathcal{O}_1:WellBore$.

Structural Indexing. Given that queries over the structural relationships of ontologies are heavily employed in our approach, we rely on the optimized structural index of LogMap [33, 39], based on the interval labelling schema techniques presented in [1].

Specifically, the structural index exploits an optimized data structure for storing directed acyclic graphs (DAGs), and it allows us to answer many entailment queries over the concept hierarchy as an index lookup operation, and hence without the need of an OWL 2 reasoner (after the initial classification of the ontology). This kind of index has demonstrated to significantly reduce the cost of answering taxonomic queries [8, 59] and disjointness relationships queries [33, 37]. A formal definition of the structural index is provided in [73], Definition 6.4, Section 6.3.2.

Hence, the concept hierarchies provided by an OWL 2 reasoner and the explicit disjointness axioms of $\mathcal{O}_1$ and $\mathcal{O}_2$ are efficiently indexed into the structural index (lines 2 and 3 in Algorithm 2).

Disjointness Addition. In order to reduce the conservativity problem to a mapping incoherence repair problem, following the notion of *assumption of disjointness*, we need to automatically add sufficient disjointness axioms into each ontology $\mathcal{O}_i$. However, the insertion of additional disjointness axioms δ may lead to unsatisfiable concepts in $\mathcal{O}_i \cup \{\delta\}$, as shown in Example 5.2.

Example 5.2. Consider the axiom β_9 from Table 1. Following the *assumption of disjointness* a very naïve algorithm would add disjointness axioms between *Borehole*, *Continuant* and *Occurrent*, which would make *Borehole* unsatisfiable.

In order to detect if each candidate disjointness axiom leads to an unsatisfiability, a non naïve algorithm requires to make an extensive use of an OWL 2 reasoner to check if there are new unsatisfiable concepts in $\mathcal{O}_i \cup \{\delta\}$. In large ontologies, however, this approach can be prohibitive.

Algorithm 3 disjointAxiomsExtension function for disjointness axioms extension**Input:** $\mathcal{P}$: propositional theory; SI : structural index; $SI_{\mathcal{U}}$: structural index of the aligned ontology**Output:** $\mathcal{P}^d$: extended propositional theory; $disj$: number of disjointness rules

```

1:  $disj \leftarrow 0$ 
2:  $\mathcal{P}^d \leftarrow \mathcal{P}$ 
3: for each  $A \rightarrow B$  in subsumptionViolations( $SI, SI_{\mathcal{U}}$ ) do ▷ As in Definition 2.6
4:   if not ( $SI.areDisj(A, B)$ ) then
5:      $\mathcal{P}^d \leftarrow \mathcal{P}^d \cup \{A \wedge B \rightarrow \perp\}$ 
6:      $SI \leftarrow SI.updateIndex(A \sqcap B \rightarrow \perp)$ 
7:      $disj \leftarrow disj + 1$ 
8:   end if
9: end for
10: return  $\langle \mathcal{P}^d, disj \rangle$ 

```

Our method, instead, exploits the propositional encoding and structural indexing of the input ontologies. Thus, checking if $\mathcal{O}_i \cup \{\delta\}$ contains unsatisfiable concepts is restricted to the Horn propositional case. The safety conditions for disjointness clauses addition to a (satisfiable) Horn propositional formula is given in Proposition 5.1, while the proof can be found in [73], Proposition 6.2, Section 6.3.2. These conditions can easily be tested with the structural index.

Proposition 5.1. Given a satisfiable Horn propositional formula $\mathcal{P}$, the addition of a (j -th) disjointness clause $A \wedge B \rightarrow \perp$ will not cause any (potential) unsatisfiability of the propositions in $\mathcal{P}$ iff: (i) neither $\mathcal{P} \models A \rightarrow B$ nor $\mathcal{P} \models B \rightarrow A$ holds, (ii) no proposition C exists such that $\mathcal{P} \models C \rightarrow A$ and $\mathcal{P} \models C \rightarrow B$.

Algorithm 2 extends the propositional formulas $\mathcal{P}_1$ and $\mathcal{P}_2$ with disjointness rules of the form $A \wedge B \rightarrow \perp$ (lines 6-7). The disjointness addition follows Proposition 5.1 and guarantees that, for every proposition A in the extended propositional formula $\mathcal{P}_i^d$ (with $i \in \{1, 2\}$), the formula $\mathcal{P}_i^d \cup \{\top \rightarrow A\}$ is satisfiable. This does not necessarily hold if the disjointness axioms are added to the OWL 2 ontologies $\mathcal{O}_1$ and $\mathcal{O}_2$, as discussed above.

Note that the addition of all possible disjointness rules may be prohibitive for large ontologies and unnecessary (see [75]). Thus, one should only add disjointness where a subsumption violation occurs, *i.e.*, adding a disjointness axiom between each pair of concepts $A, B \in \mathcal{O}_i$ (with $i \in \{1, 2\}$) such that $A \sqsubseteq B \in \text{subViol}(\mathcal{O}_i, \mathcal{O}_{\mathcal{O}_1, \mathcal{O}_2}^{\mathcal{M}})$, as in Definition 2.6. Algorithm 3 implements this idea for the Horn propositional case and extensively exploits the structural index to identify the subsumption violations (line 3 of Algorithm 3). This algorithm requires as input the structural index of the integrated ontology, and thus its classification with an OWL 2 reasoner (line 5 in Algorithm 2). The classification time of the integrated ontology is known to be typically much higher than that of the input ontologies individually [34]. However, this was not a bottleneck in our experiments, as shown in Section 7.2.

Mapping Repair. The step 8 of Algorithm 2 uses the mapping (incoherence) repair algorithm of LogMap, for the extended Horn propositional formulas $\mathcal{P}_1^d$ and $\mathcal{P}_2^d$, and the input mappings $\mathcal{M}$. The mapping repair process exploits the Dowling-Gallier (D&G) algorithm [14, 23] for propositional Horn satisfiability (refer to [73], Section 6.3, for more details) and checks, for every proposition A of a given formula $\mathcal{P}$, the satisfiability of the propositional formula $\mathcal{P}_A = \mathcal{P} \cup \{\top \rightarrow A\}$. Satisfiability of $\mathcal{P}_A$ is checked in worst-case linear time in the size of $\mathcal{P}_A$, and the number of D&G calls is also linear in the number of propositions in $\mathcal{P}$. In case of unsatisfiability, the variant of the algorithm

implemented in LogMap also allows us to record *conflicting* mappings involved in the unsatisfiability, which will be considered for the subsequent repair process.¹⁷

The repair is performed by removing some of the identified conflicting mappings. To this aim, the algorithm selects an (approximation) of the minimal hitting set (w.r.t. total mapping confidence) between the set of conflicting mappings, for the different unsatisfiabilities that are detected. In case of multiple options, the mapping confidence will be used as a differentiating factor.¹⁸ Example 5.3 provides an example of mapping repair for solving a subsumption violation.

Example 5.3. Consider the propositional encoding $\mathcal{P}_1$ and $\mathcal{P}_2$ of the axioms of Table 1 and the mappings $\mathcal{M}$ in Table 2, seen as propositional rules. $\mathcal{P}_1^d$ and $\mathcal{P}_2^d$ have been created by adding disjointness rules to $\mathcal{P}_1$ and $\mathcal{P}_2$, according to Algorithm 3. For example, $\mathcal{P}_2^d$ includes the rule $\psi = \mathcal{O}_2:Well \wedge \mathcal{O}_2:Borehole \rightarrow \text{false}$. The mapping repair algorithm identifies the propositional theory $\mathcal{P}_1^d \cup \mathcal{P}_2^d \cup \mathcal{M} \cup \{\text{true} \rightarrow \mathcal{O}_1:ExplorationWellbore\}$ as unsatisfiable. This is due to the combination of the mappings m_3 and m_4 , the propositional projection of axioms β_1 and β_2 , and the rule ψ . The mapping repair algorithm also identifies m_3 and m_4 as the cause of the unsatisfiability, and discards m_3 , since its confidence is smaller than that of m_4 (see Table 2).

6. Properties of the Repair Methods and Combined Approach

We have presented two methods for the detection and correction of violations of the conservativity principle. The repair of subsumption conservativity violations (which are reduced to a consistency violations) is based on the Dowling-Gallier algorithm for propositional Horn satisfiability; whereas equivalence conservativity violations are addressed using a combination of graph theory and logic programming.

Soundness and (In)completeness. Both methods are sound (the violations that are detected are indeed violations if considering the full expressiveness of the input ontologies), but incomplete, since the used approximate projections of the input ontologies (*i.e.*, Horn propositional and graph encodings) may lead to some violations being missed.

Algorithm 2 computes a repair $\mathcal{R}^\approx$ such that $\mathcal{M}' = \mathcal{M} \setminus \mathcal{R}^\approx$ is coherent with respect to $\mathcal{P}_1^d$ and $\mathcal{P}_2^d$ (according to the propositional case of Definition 2.5). Furthermore, the propositional theory $\mathcal{P}_1 \cup \mathcal{P}_2 \cup \mathcal{M}'$ does not contain any subsumption violation with respect to $\mathcal{P}_1$ and $\mathcal{P}_2$ (according to the propositional case of Definition 2.6). However, our encoding is incomplete, and we cannot guarantee that $\mathcal{O}_1 \cup \mathcal{O}_2 \cup \mathcal{M}'$ does not contain any subsumption conservativity violations w.r.t. $\mathcal{O}_1$ and $\mathcal{O}_2$.

Analogously, Algorithm 1 computes a diagnosis (*i.e.*, minimal mapping repair) Δ such that $\mathcal{M}' = \mathcal{M} \setminus \Delta$ does not lead to equivalence violations with respect to the graph encoding of the aligned ontology presented in Definition 4.1. However, the considered encoding is incomplete since only the classification of the input ontologies in isolation is considered. Thus, as for the subsumption violations, we cannot guarantee that $\mathcal{O}_1 \cup \mathcal{O}_2 \cup \mathcal{M}'$ does not contain any equivalence violation.

¹⁷ Note that, as for the case of EqRepair, we do not compute the classification of $\mathcal{O}_1 \cup \mathcal{O}_2 \cup \mathcal{M}$ since this will prevent our extension of D&G to identify and record the mappings involved in an unsatisfiability (*i.e.*, a subsumption violation).

¹⁸ In scenarios where the confidence of the mapping is missing (*e.g.*, in reference or manually created mapping sets) or unreliable, our mapping repair technique computes fresh confidence values based on the locality principle [36].

Algorithm 4 ConsRepair, a multi-strategy approach to repair conservativity violations**Input:** $\mathcal{O}_1, \mathcal{O}_2$: input ontologies; $\mathcal{M}$: input (coherent) mappings; rs : repair strategy;**Output:** $\mathcal{R}^\approx$: approximate repair; $disj$: number of disjointness rules

```

1:  $disj \leftarrow 0$ 
2: switch  $rs$  do
3:   case equivalence repair
4:      $\mathcal{R}^\approx \leftarrow \text{EqRepair}(\mathcal{O}_1, \mathcal{O}_2, \mathcal{M})$  ▷ Algorithm 1
5:   case subsumption repair
6:      $\langle \mathcal{R}^\approx, disj \rangle \leftarrow \text{SubRepair}(\mathcal{O}_1, \mathcal{O}_2, \mathcal{M})$  ▷ Algorithm 2
7:   case equivalence and subsumption repair
8:      $\mathcal{R}_1^\approx \leftarrow \text{EqRepair}(\mathcal{O}_1, \mathcal{O}_2, \mathcal{M})$ 
9:      $\langle \mathcal{R}_2^\approx, disj \rangle \leftarrow \text{SubRepair}(\mathcal{O}_1, \mathcal{O}_2, \mathcal{M} \setminus \mathcal{R}_1^\approx)$ 
10:     $\mathcal{R}^\approx \leftarrow \mathcal{R}_1^\approx \cup \mathcal{R}_2^\approx$ 
11:   case subsumption and equivalence repair
12:     $\langle \mathcal{R}_1^\approx, disj \rangle \leftarrow \text{SubRepair}(\mathcal{O}_1, \mathcal{O}_2, \mathcal{M})$ 
13:     $\mathcal{R}_2^\approx \leftarrow \text{EqRepair}(\mathcal{O}_1, \mathcal{O}_2, \mathcal{M} \setminus \mathcal{R}_1^\approx)$ 
14:     $\mathcal{R}^\approx \leftarrow \mathcal{R}_1^\approx \cup \mathcal{R}_2^\approx$ 
15: return  $\langle \mathcal{R}^\approx, disj \rangle$ 

```

Nevertheless, incompleteness is mitigated thanks to the classification of the input ontologies using full reasoning (see Definition 4.1 and the propositional encoding in Section 5). Furthermore, our evaluation suggests that the number of remaining violations after repair is typically small (see Section 7.2).

Combined Algorithm. Despite the differences between subsumption and equivalence violations, a mutual influence between them exists. An example is given in Table 3 of Section 3, where subsumption violations are caused by the existence of equivalence violations in the same aligned ontology. Nonetheless, subsumption violations can also be responsible for the existence of equivalence violations.

It is evident, from the aforementioned example, that the application of a repair step targeting a particular violation kind can solve as a side-effect violations of other kinds.

In Algorithm 4 we provide a multi-strategy algorithm to repair conservativity violations. The algorithm allows combining the approaches targeting subsumption and equivalence violations. We expect that the combination of the repair strategies will impact the performance in terms of required time and number of solved violations. Furthermore, we also expect the order of the strategies will influence the results.

Algorithm 4, as SubRepair and EqRepair algorithms, expects as input two ontologies $\mathcal{O}_1$ and $\mathcal{O}_2$, and a coherent alignment $\mathcal{M}$. Additionally it also expects the required repair strategy rs (e.g., subsumption repair first and then equivalence repair). As output the algorithm provides the (approximate) repair $\mathcal{R}^\approx$ and the number of added disjointness $disj$ in the SubRepair algorithm. When SubRepair is not executed $disj = 0$. In the case of using a combined strategy, the repaired mappings by the SubRepair (resp. EqRepair) algorithm will be given as input of the EqRepair (resp. SubRepair) algorithm, see lines 12 and 13 (resp. 8 and 9) in Algorithm 4.

7. Experimental Evaluation

In this section, we evaluate the feasibility of using our combined approach to correct conservativity violations in practice.^{19,20} To this end we have conducted the evaluation in Algorithm 5 over the ontologies and mapping sets (*i.e.*, reference alignments and alignments computed by participating systems) of the 2012–2014 campaigns of the Ontology Alignment Evaluation Initiative (OAEI):

- i *Anatomy*: the Anatomy dataset involves the Adult Mouse Anatomy (*MO*) ontology and a fragment of the *NCI* ontology describing human anatomy. The reference alignment has been manually curated [86]. From 2012 to 2014 we have gathered 47 different alignments computed by participating systems.
- ii *Conference*: this dataset uses a collection of 16 ontologies from the domain of academic conferences [83]. Currently, there are 21 manually created mapping sets among 7 of the ontologies. From 2012 to 2014 we have gathered 1,104 different alignments computed by participating systems.
- iii *LargeBio*: this dataset includes the biomedical ontologies *FMA*, *NCI* and (a fragment of) *SNOMED CT*, and three reference mapping sets based on the *UMLS* [5]. From 2012 to 2014 we have gathered 122 alignments computed by participating systems.
- iv *Library*: this OAEI dataset includes the real-word thesauri *STW* and *TheSoz* from the social sciences. The reference mappings have been manually validated. From 2012 to 2014 we have gathered 32 alignments computed by participating systems.

We have run the evaluation algorithm for (i) each of the OAEI tasks (*i.e.*, pair of ontologies) described above and (ii) each of the reference and computed alignments available in each OAEI task, which resulted in 1,331 executions and 5,324 calls (1,331 x 4 repair strategies) to the repair method (line 9 in Algorithm 5).

In the conducted evaluation we have used the OWL 2 reasoner *HermiT* [25] to classify the input and aligned ontologies. In a few cases *HermiT* failed in classifying the aligned ontology and we used the OWL 2 EL reasoner *ELK* [43] in order to provide an approximation of the classification. The use of this approximation had a minor impact in the overall results.

Section 7.1 briefly comments on the steps followed by Algorithm 5. In Section 7.2 we present and discuss the results of our repair methods in terms of computation times, computed repairs and corrected violations. Note that we have grouped the results according to the origin of the alignments: reference and computed by participating systems. Section 7.3 analyses and compares the impact of the conservativity repair strategies on the alignment quality in terms of precision, recall, and f-measure.

¹⁹ The complete source code of the proposed algorithms and the performed experiments is available at <https://github.com/asolimando/logmap-conservativity/>

²⁰ The test environment consisted of a desktop computer equipped with 32GB DDR3 RAM at 1333MHz, and an AMD Fusion FX 4350 (quad-core, each running at 4.2GHz) as CPU. The dataset is stored on a 128GB SSD, where the operating system *Ubuntu* (12.04, 64-bit version) is also installed. Our prototype can run with less than 8GB for the majority of the considered tests, we however allocate 26GB of RAM for the JVM in order to minimize the influence of the garbage collector on the recorded temporal measurements.

Algorithm 5 Conducted evaluation

Input: $\mathcal{O}_1, \mathcal{O}_2$: input ontologies; $\mathcal{M}$: alignment between $\mathcal{O}_1$ and $\mathcal{O}_2$; RS: repair strategies;

Pre-processing steps:

- 1: $\langle \mathcal{O}_1, \mathcal{O}_2 \rangle \leftarrow \text{ModuleExtraction}(\mathcal{O}_1, \mathcal{O}_2, \mathcal{M})$
- 2: $\mathcal{M} \leftarrow \text{ConsistencyRepair}(\mathcal{O}_1, \mathcal{O}_2, \mathcal{M})$
- 3: $\mathcal{O}^{\mathcal{M}} \leftarrow \mathcal{O}_1 \cup \mathcal{O}_2 \cup \mathcal{M}$ ▷ Definition 2.3
- 4: Store $|\text{Sig}(\mathcal{O}_1)|$, $|\text{Sig}(\mathcal{O}_2)|$, and $|\mathcal{M}|$
- 5: $\text{subViol}_i \leftarrow |\text{subViol}(\mathcal{O}_1, \mathcal{O}^{\mathcal{M}})| + |\text{subViol}(\mathcal{O}_2, \mathcal{O}^{\mathcal{M}})|$ ▷ Subsumption violations, Definition 2.6
- 6: $\text{eqViol}_i \leftarrow |\text{eqViol}(\mathcal{O}_1, \mathcal{O}^{\mathcal{M}})| + |\text{eqViol}(\mathcal{O}_2, \mathcal{O}^{\mathcal{M}})|$ ▷ Equivalence violations, Definition 2.6
- 7: $\text{diff}_i^{\approx} \leftarrow |\text{diff}_{\text{Sig}(\mathcal{O}_1)}^{\approx}(\mathcal{O}_1, \mathcal{O}^{\mathcal{M}})| + |\text{diff}_{\text{Sig}(\mathcal{O}_2)}^{\approx}(\mathcal{O}_2, \mathcal{O}^{\mathcal{M}})|$ ▷ General notion, Definition 2.4
- 8: **for each** rs **in** RS **do** ▷ Repair strategies as in Algorithm 4

Compute repair (Algorithm 4):

 - 9: $\langle \mathcal{R}^{\approx}, \text{disj} \rangle \leftarrow \text{ConsRepair}(\mathcal{O}_1, \mathcal{O}_2, \mathcal{M}, rs)$ ▷ Keep times to compute repair t_r and disjointness rules t_d
 - 10: $\mathcal{M}' \leftarrow \mathcal{M} \setminus \mathcal{R}^{\approx}$
 - 11: $\mathcal{O}^{\mathcal{M}'} \leftarrow \mathcal{O}_1 \cup \mathcal{O}_2 \cup \mathcal{M}'$
 - 12: $\text{subViol}_r \leftarrow |\text{subViol}(\mathcal{O}_1, \mathcal{O}^{\mathcal{M}'})| + |\text{subViol}(\mathcal{O}_2, \mathcal{O}^{\mathcal{M}'})|$
 - 13: $\text{eqViol}_r \leftarrow |\text{eqViol}(\mathcal{O}_1, \mathcal{O}^{\mathcal{M}'})| + |\text{eqViol}(\mathcal{O}_2, \mathcal{O}^{\mathcal{M}'})|$
 - 14: $\text{diff}_r^{\approx} \leftarrow |\text{diff}_{\text{Sig}(\mathcal{O}_1)}^{\approx}(\mathcal{O}_1, \mathcal{O}^{\mathcal{M}'})| + |\text{diff}_{\text{Sig}(\mathcal{O}_2)}^{\approx}(\mathcal{O}_2, \mathcal{O}^{\mathcal{M}'})|$
- 15: **end for**

7.1. Evaluation Algorithm Steps

Algorithm 5 expects as input two ontologies $\mathcal{O}_1$ and $\mathcal{O}_2$, a (not necessarily coherent) alignment $\mathcal{M}$, and the set of supported repair strategies RS (as in Algorithm 4). Note that the evaluation algorithm includes two pre-processing steps of the input mappings and ontologies (lines 1 and 2 of Algorithm 5).

Module Extraction. Modules are a general technique for ensuring scalability of ontology based algorithms. The main intuition is that a (logic-based) module for a given ontology $\mathcal{O}$ w.r.t. a seed signature Σ , is a subset of $\mathcal{O}$ that preserves the entailment relation over axioms expressed using Σ . In our approach, in order to reduce the size of the problem, we (optionally) extract two modules [9, 10], one for each input ontology, using the entities involved in the input mappings $\mathcal{M}$ as seed signatures. When no confusion arises, we directly refer to these modules as the input ontologies.

Consistency repair. As already mentioned, our (combined) approach expects to work on coherent alignments. In order to meet this requirement, we perform a preliminary consistency repair as a two-step process. First, the consistency repair facility of LogMap [39] is used to repair the aligned ontology. Then, any unsolved incoherence is detected through an OWL 2 reasoner, and solved by computing a single justification [28], and removing the mapping with least confidence appearing in it. This last step is iteratively applied until no more incoherences are present in the aligned ontology. This consistency repair process is sound and complete, but it does not guarantee the optimal minimization of the total confidence of the removed mappings. However, for the purposes of our evaluation, this two-step process is sufficient.

Problem characteristics. Line 4 keeps the size of the problem (*i.e.*, number of entities of the input ontologies and size of the input mapping). The number of initial violations prior repair are also stored: subsumption violations (line 5), equivalence violations (line

6) and conservativity violations following the general notion (line 7), which fully relies on the approximation of the deductive difference (Definition 2.4).

Conservativity repair. Algorithm 5 computes a repair for each of the four supported strategies keeping the times to compute the repair (t_r) and to add the disjointness rules (t_d) in the relevant cases (line 9). Once the repair has been computed, it also calculates the violations that have been missed (lines 10-14). Although our approach is incomplete the number of remaining violations is typically negligible compared to the initial number of violations reported in lines 5-7.

7.2. Evaluation on Computed and Reference Alignments

The result of applying the evaluation algorithm to the ontologies and reference alignments of the OAEI campaign is presented in Figure 7.²¹ Figure 8 reports the same evaluation over the computed alignments by participating systems in the OAEI 2012–2014 campaigns. Tables 7a and 8a show the characteristics of the problem (*i.e.*, average size of input ontologies and mappings, and number of violations prior repair). The last column ($\#\mathcal{M}$) represents the number of matching tasks in each of the tracks. For example, there are 122 different alignments in the *largebio* track computed by participating systems with an average size of 11,789 mappings per alignment. The results compiling the output of several repairs (*i.e.*, over several alignments) are presented in the tables with the format *mean (standard deviation)*. For example, for the *conference* track (Table 7a), since $\#\mathcal{M}=21$ (*i.e.*, there are 21 alignments), we report the average size of (initial) subsumption violations (2.19) and the standard deviation of their size (3.5).

Next, we first analyse the results of the two basic repair strategies in isolation (*i.e.*, EqRepair and SubRepair), and then we compare with the results of the combined approaches (*i.e.*, EqRepair first and then SubRepair, and vice versa).

Equivalence. The experimental results considering EqRepair algorithm, can be summarized as follows:

- (i) The sum of the detection and repair time of EqRepair is very low due to the linear cost of the detection technique and the efficient parallelization of the diagnosis computation.
- (ii) The computed repairs are typically of limited size (less than 10%), but can reach a significant portion of the the original alignment. For example, in the *library* track (Table 8b) it removes on average 22% of the mappings. This is due to the large number of violations led by the alignments computed by participating systems.
- (iii) The equivalence violations are completely removed for *anatomy* and *library* tracks, and practically removed for *largebio* and *conference* compared to the initial number of equivalence violations (see Tables 7b and 8b).
- (iv) The removal of equivalence violations, as expected, has also an important impact w.r.t. the remaining subsumption violations. Specially when dealing with mapping sets computed by participating systems in the OAEI's *library* track (see Table 8b).

This set of experiments confirms the effectiveness and efficiency of the detection and repair algorithms for equivalence violations.

²¹ Note that the reference mappings of the OAEI campaign are already coherent w.r.t. the test case ontologies and thus the consistency repair step was not necessary.

OAEI track	Problem Size			Initial Violations			# $\mathcal{M}$
	$ \text{Sig}(\mathcal{O}_1) $	$ \text{Sig}(\mathcal{O}_2) $	$ \mathcal{M} $	subViol _i	diff _i [≈]	eqViol _i	
anatomy	2,747	3,306	3,032	1,321	1,335	26	1
conference	110(31)	121(44)	29(10)	2.19(3.5)	2.24(3.48)	0.1(0.3)	21
largebio	21,693(25,739)	14,674(8,811)	19,813(15,457)	234,111(284,869)	242,517(297,275)	2,536(3,096)	3
library	6,575	8,376	6,322	42,045	42,872	895	1

(a) Problem characteristics.

OAEI track	Solution Size	Time	Remaining Violations		
	$ \mathcal{R}^{\approx} $	t_r (s)	subViol _r	diff _r [≈]	eqViol _r
anatomy	25	0.61	1,259	1,271	0
conference	0.1(0.3)	0.05(0.05)	2.05(3.4)	2.05(3.4)	0
largebio	1,801(1,927)	15(5.28)	164,377(187,554)	168,819(194,184)	18(32)
library	627	2.58	30,006	30,408	0

(b) Equivalence repair in isolation.

OAEI track	Solution Size		Times		Remaining Violations		
	#disj	$ \mathcal{R}^{\approx} $	t_d (s)	t_r (s)	subViol _r	diff _r [≈]	eqViol _r
anatomy	1,321	317	0.68	0.35	0	3	1
conference	2.19(3.5)	1.05(1.47)	0.05(0.05)	0	0	0.05(0.22)	0.05(0.22)
largebio	234,111(284,869)	6,625(5,749)	28(22)	164(243)	170(213)	688(944)	123(175)
library	42,045	2,183	2.73	9.31	0	38	10

(c) Subsumption repair in isolation.

OAEI track	Solution Size		Times		Remaining Violations		
	#disj	$ \mathcal{R}^{\approx} $	t_d (s)	t_r (s)	subViol _r	diff _r [≈]	eqViol _r
anatomy	1,259	325	0.68	0.53	0	2	0
conference	2.05(3.4)	1.05(1.32)	0.05(0.05)	0.01(0.01)	0	0	0
largebio	164,377(187,554)	6,683(5,912)	27(20)	114(162)	36(57)	121(68)	0.33(0.58)
library	30,006	2,140	2.45	8.77	0	6	0

(d) Equivalence repair first and then subsumption repair (Eq.→Sub).

OAEI track	Solution Size		Times		Remaining Violations		
	#disj	$ \mathcal{R}^{\approx} $	t_d (s)	t_r (s)	subViol _r	diff _r [≈]	eqViol _r
anatomy	1,321	318	0.68	0.47	0	2	0
conference	2.19(3.5)	1.1(1.45)	0.05(0.05)	0.01(0.01)	0	0	0
largebio	234,111(284,869)	6,730(5,897)	28(22)	165(244)	48(51)	158(141)	1(1.73)
library	42,045	2,193	2.73	9.73	0	9	0

(e) Subsumption repair first and then equivalence repair (Sub.→Eq.).

Fig. 7. Results for OAEI reference alignments, grouped by track. Cells compiling the results of several repairs are presented with the format “mean (std deviation)”.

Subsumption. The experimental results on the behaviour of SubRepair algorithm, can be summarized as follows:

- (i) Despite the large number of subsumption violations, the required time t_d to detect violations and add disjointness rules (*i.e.*, cost of Algorithm 3) is very low (see Tables 7c and 8c). The results also show that t_d is directly influenced by the size of the input ontologies, more than by the number of violations.
- (ii) Repair times t_r are typically small and they do not represent a bottleneck in spite of the large number of added disjointness rules. An exception is represented by the *library* track (Table 8c), where an average repair time of 33 minutes is required. However, the runtime is reasonable considering the impressive average number of violations (*i.e.*, 5M).
- (iii) The computed repairs $\mathcal{R}^{\approx}$ can be aggressive and the average mapping removal

OAEI track	Problem Size			Initial Violations			# $\mathcal{M}$
	$ \text{Sig}(\mathcal{O}_1) $	$ \text{Sig}(\mathcal{O}_2) $	$ \mathcal{M} $	subViol_i	$\text{diff}_i^{\approx}$	eqViol_i	
anatomy	2,747	3,306	2,607(816)	6,562(12,844)	6,619(12,911)	137(340)	47
conference	110(30)	121(43)	33(35)	17(86)	17(87)	1.27(7.74)	1,104
largebio	18,340(19,857)	13,450(6,995)	11,789(8,729)	154,566(237,174)	159,711(245,746)	716(1,100)	122
library	6,575	8,376	6,363(3,168)	5,121,627(8*10 ⁶)	5,153,135(8*10 ⁶)	516,289(10 ⁶)	32

(a) Problem characteristics.

OAEI track	Solution Size	Time	Remaining Violations		
	$ \mathcal{R}^{\approx} $	t_r (s)	subViol_r	$\text{diff}_r^{\approx}$	eqViol_r
anatomy	82(190)	4.24(14)	4,011(7,888)	4,030(7,913)	0
conference	0.85(5.83)	0.26(3.48)	11(42)	11(43)	0.05(1.16)
largebio	453(661)	13(11)	110,323(153,602)	113,571(158,770)	18(49)
library	1,401(1,866)	24(37)	90,154(53,457)	91,996(54,483)	0

(b) Equivalence repair in isolation.

OAEI track	Solution Size		Times		Remaining Violations		
	#disj	$ \mathcal{R}^{\approx} $	t_d (s)	t_r (s)	subViol_r	$\text{diff}_r^{\approx}$	eqViol_r
anatomy	6,562(12,844)	396(413)	0.77(0.4)	1.22(2.6)	0	2.23(1.99)	0.47(1.02)
conference	17(86)	3.4(10)	0.17(1.45)	0.01(0.09)	0.15(1.33)	0.18(1.34)	0.03(0.2)
largebio	154,566(237,174)	2,897(2,847)	25(18)	134(286)	79(151)	234(348)	35(61)
library	5,121,627(8*10 ⁶)	2,865(2,196)	48(75)	1,966(3,690)	0	67(54)	16(14)

(c) Subsumption repair in isolation.

OAEI track	Solution Size		Times		Remaining Violations		
	#disj	$ \mathcal{R}^{\approx} $	t_d (s)	t_r (s)	subViol_r	$\text{diff}_r^{\approx}$	eqViol_r
anatomy	4,011(7,888)	384(400)	0.76(0.4)	4.39(14)	0	1.6(0.85)	0
conference	11(42)	3.43(10)	0.16(1.41)	0.22(3.48)	0.16(1.24)	0.16(1.24)	0
largebio	110,323(153,602)	2,851(2,823)	24(17)	87(145)	52(130)	104(140)	0.43(1.45)
library	90,154(53,457)	2,896(2,329)	39(62)	42(42)	0	3.81(4)	0

(d) Equivalence repair first and then subsumption repair (Eq. $\rightarrow$ Sub).

OAEI track	Solution Size		Times		Remaining Violations		
	#disj	$ \mathcal{R}^{\approx} $	t_d (s)	t_r (s)	subViol_r	$\text{diff}_r^{\approx}$	eqViol_r
anatomy	6,562(12,844)	397(414)	0.77(0.4)	1.31(2.6)	0	1.51(0.91)	0
conference	17(86)	3.43(10)	0.17(1.45)	0.01(0.09)	0.15(1.33)	0.15(1.33)	0
largebio	154,566(237,174)	2,926(2,890)	25(18)	135(287)	53(129)	107(144)	0.51(1.48)
library	5,121,627(8*10 ⁶)	2,881(2,204)	48(75)	1,967(3,690)	0	15(10)	0

(e) Subsumption repair first and then equivalence repair (Sub. $\rightarrow$ Eq.).Fig. 8. Results for OAEI 2012-2014 computed alignments, grouped by track. Cells compiling the results of several repairs are presented with the format “*mean (std deviation)*”.

ranges from 3% (*conference* track, Table 7c) up to 45% of the original alignment (*library* track, Table 8c). From Tables 7c and 8c, the expected positive correlation between the repair size and the number of detected violations (*i.e.*, added disjointness rules #disj) clearly emerges.

- (iv) Subsumption violations are completely removed in the *anatomy* and *library* cases, and almost completely removed in the *conference* and *largebio* tracks, with less than 0.1% (on average) of unsolved violations. As a side effect equivalence violations are also reduced to a minimum. For example, for one of the computed alignments between *iasted* and *sigkdd* ontologies in the *conference* track we fail to detect and repair the conservativity violation $iasted:Hotel_fee \sqsubseteq iasted:Reg_fee$. The justifications of this violation reveal that an inverse property axiom and an existential restriction are behind this novel entailment. These OWL 2 constructors,

however, fall outside the Horn propositional and graph projections of the input ontologies currently implemented in our techniques.

- (v) The number of missed violations is only slightly higher when considering the general notion of the conservativity principle (see $\text{diff}_r^\approx$ columns in Tables 7c and 8c), which suggests that our (approximate) variant based on the assumption of disjointness is suitable in practice. Furthermore, the number of unsolved violations using this notion is negligible (less than 0.25% on average).

This set of experiments basically confirms the effectiveness and efficiency of the detection and repair algorithms for subsumption violations. This evaluation also shows that the violation detection algorithm is more heavily influenced by the size of the involved ontologies and alignments, while the time required by the repair algorithm is strongly correlated with the number of violations to repair.

Combined. The experimental results on the behaviour of the different flavours of our combined approaches can be summarized as follows:

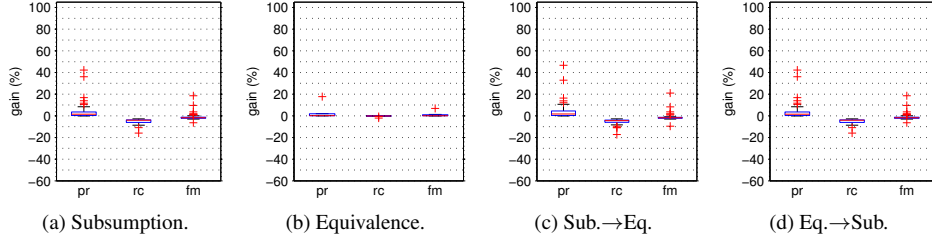
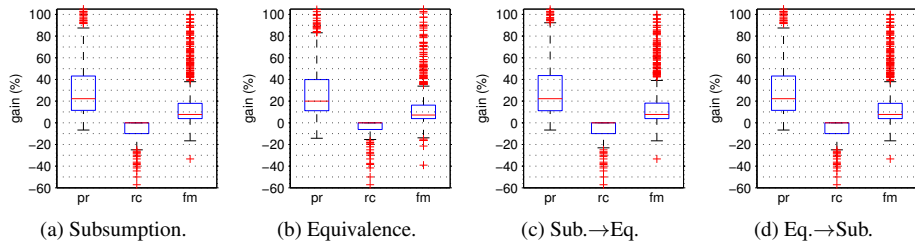
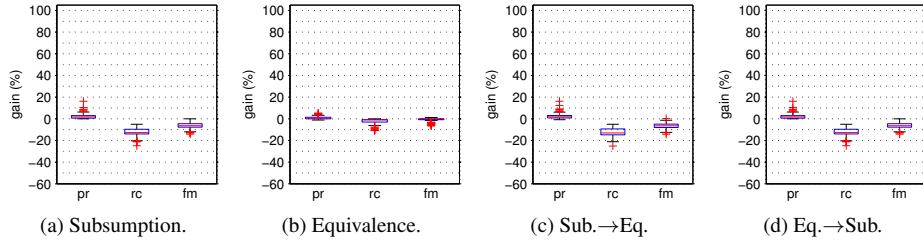
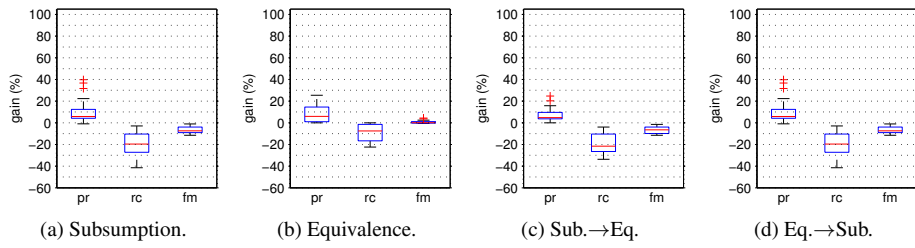
- (i) In Tables 7d, 7e, 8d, and 8e we can see that the number of unsolved violations, as expected, is lowered w.r.t. applying one of the repair methods in isolation. In *anatomy* and *conference* the results are slightly better when applying SubRepair algorithm first, while in *largebio* and *library* the unsolved violations are smaller when applying EqRepair algorithm first; nevertheless the differences are negligible.
- (ii) The detection and disjointness addition time t_d is comparable in both combined variants, and it is not therefore influenced by the repair application order. However, the number of required disjointness rules $\#disj$ is significantly reduced when applying EqRepair algorithm first.
- (iii) Total repair times t_r when applying SubRepair first are basically the same as the times when applying SubRepair in isolation. However, the repair times when applying EqRepair first for the *largebio* and *library* tracks are significantly reduced, specially for the case of *library*. This is due to the fact that the SubRepair algorithm, in presence of SCCs, generates a large number of almost equivalent repair plans, that pose a problem to the plan selection method. The solution presented in EqRepair is more efficient in these cases since it focuses on the problematic SCCs.
- (iv) The computed repairs $\mathcal{R}^\approx$ are of comparable size, on average smaller when equivalence repair is applied first. An exception is, however, represented by the alignments computed by participating systems in the *library* track.

This set of experiments confirms the increased effectiveness, with a small loss in efficiency, of the combined repair approaches, w.r.t. the single repair methods.

7.3. Repair Effects on Alignment Quality

Figures 9–12 show the average impact (in terms of precision, recall and f-measure) of applying a conservativity repair over the computed alignments by participating systems in the OAEI 2012–2014 campaigns. The results are grouped by track.

The impact of alignment repair is computed as the percentual of gain (resp. loss for negative values) for each measure computed for a repaired alignment, compared to the same measure computed for the original alignment. The measures are computed w.r.t. the OAEI reference alignments. Additionally, we filter any alignment without conservativity violations, because the empty repair always implies a void gain.

Fig. 9. Repair impact for *anatomy* track.Fig. 10. Repair impact for *conference* track.Fig. 11. Repair impact for *largebio* track.Fig. 12. Repair impact for *library* track.

Such impact of the conservativity repair is presented in Figures 9–12 for each flavour of our repair algorithm, where results over individual alignments are grouped by track. The boxplots presented in such figures are characterized by a “box” (*i.e.*, the range in between the first and third percentile), the “median” for the results, and the “whiskers” which allows to identify outliers (depicted as red crosses). Consider, for instance, Figure 10a, where the box ranges approximately from 6% to 42%, having a median value of 21% and lower (resp. upper) whisker at approximately -9% (resp. 89%).

Figures 9–12 confirm the general effects of a repair algorithm (*i.e.*, an increase of precision and a decrease of recall). The results achieved by the two versions of the combined repair algorithm are totally comparable. The gain in precision results in an increase of the resulting f-measure for *conference* track, and a slight decrease for *anatomy* track. However, for *largebio* and *library* tracks, the loss in f-measure is higher, and lies in the range $[0, 10]\%$.

From the results it is evident that the effect of the combined repair algorithm is in line with that of the repair algorithm addressing exclusively subsumption violations, in terms of performance w.r.t. a reference alignment. The correlation between the repair size and its impact is again evident.

We also remind that the measures are computed against unrepaired reference alignments, while a more appropriate comparison would be against reference alignments where all the true positive conservativity violations are repaired. For instance, having the highest gain in f-measure for the *anatomy* and *conference* tracks, seems to confirm our hypothesis because, in these cases, the reduced size of the input ontologies and reference alignments allows for a more effective revision, which also limits the conservativity violations (consider the reduced number of violations in Table 8a for the *conference* track), that are more frequent in the other tracks (considering of course also the difference in the size of the alignments). Finally, it is not surprising that the worst result in terms of the effect on the f-measure is obtained for the *library* track, for which the extremely high number of conservativity violations (see Table 8a) suggests the need for a more accurate revision of the reference alignment and input ontologies.

8. Related Work

This section provides a general overview of the state of the art on ontology and ontology alignments debugging techniques, the related work on the *Weighted Feedback Edge Set* problem, and the notion of *assumption of disjointness*.

Conservativity Principle, Ontology vs Alignment Repair. Once a violation of the conservativity principle is detected, there are different approaches to correct it by modifying the aligned ontology (*i.e.*, targeting the input ontologies or the alignment).

The first approach is to consider as problematic all the violations and to correct them using classic ontology repair strategies on the aligned ontology, *i.e.*, to compute a hitting set over all the justifications, as discussed at the end of Section 2.3. The repair computed in this way may remove elements either from the ontologies, or from the alignment, or a mixture of the two.

The second one is a family of approaches from Lambrix *et al.* [48, 50], which targets violations that can be considered as false positives, for which the problem source is considered to stem from the incompleteness of the input ontologies. The correction strategy aims at adding to the input ontologies a minimal set of axioms, so that the input ontologies (in isolation) can entail the novel axiom (solving, in this way, the violation).

Lambrix *et al.* also propose a unified approach [31, 49]. For each detected violation

different repair plans are generated, and they are ranked w.r.t. their informativity (*i.e.*, considering the number of solved depending violations) and minimality (based on the number of inserted/deleted axioms). The user can classify the violations as false or true positives, and can in this way influence the plan ranking. Our heuristic could be conceived as an alternative and additional plan in this multistrategy approach, that supports both automatic and assisted repair.

The approaches in [31, 49, 50] refer to the same (simplified) context addressed here for equivalence violations, namely taxonomical projection of expressive DLs, claiming that this is one of the most used features in semantic-enabled applications.

Our proposal is orthogonal to the others discussed in this section, and could be an alternative inside a multistrategy approach. However, there are settings for which the only repair target is the alignment, such as Multi-Agent Systems, where the private knowledge of each agent is encoded as a OWL ontology and cannot be modified in order to agree on a common alignment [40, 41, 53, 61, 62]. In a more general ontology matching setting, it may also be the case that one (or both) of the input ontologies should not be modified. For example, the LUCADA ontology—a medium-sized OWL ontology that describes the domain of Lung Cancer according to the specifications of the National Health Service (NHS) [70]—was integrated with SNOMED CT in order to facilitate interoperability with other applications within the NHS. In this use case, SNOMED CT was considered as immutable (during the matching process) since it is the reference ontology across NHS’s information systems.

State-of-the-art ontology alignment repair systems, such as ALCOMO [54], AML [66], ASMOV [32], Lily [85], LogMap [33], and YAM++ [60], typically consider the input ontologies as immutable and their repair techniques focus on the mappings.²² Nevertheless, Pesquita *et al.* [7] question the automatic generation of repairs and suggest to update the ontologies when required.

Equivalence and Subsumption Violations. Ontology alignment violations are discussed in [4], where sanity checks for ontology mappings are proposed, together with general recommendations about best practices in producing ontology mappings (*e.g.*, use of URI for identifying ontologies and matched elements). Sanity checks 6 and 7 are particularly meaningful for our analysis. Check 6 forbids the entailment of novel equivalences, among entities of the same input ontology, caused by the presence of multiple mappings sharing a common entity. This check is analyzed in the context of an exclusive use of “ $\equiv$ ” semantic relation for an alignment between an ontology of mouse anatomy and the corresponding one related to human beings (*NCI ontology*²³). In none of the 39 detected cases, the target entities were stated equivalent in the input ontology. In a later release of such ontology, 15 entities were merged, while 18 were judged as not equivalent by domain experts (*NCI ontology curators*). This suggests that it is not possible to exclude that (new) correct equivalences may be derived using ontology matching. However, it does not seem to be the only case. Although discovering new relations between entities of the same input ontology could be very useful, their indiscriminate use, in absence of a manual exploratory phase, is only advisable for contexts in which correctness is not mandatory (*e.g.*, information retrieval), while for data integration and query answering tasks, a conservative approach seems more appropriate [78]. Check 7, instead, refers to subsumption violations, for which similar considerations apply.

Another relevant approach is the one proposed by Arnold *et al.* [2]. Despite not be-

²² The interested reader please refer to Section 4.6.4 in [73] for an overview of these approaches.

²³ <http://www.cancer.gov/cancertopics/cancerlibrary/terminologyresources>

ing directly linked to the conservativity principle and to ontology debugging in general, the proposed algorithm could mitigate the presence of violations by refining the alignment produced by ontology matchers. The proposed method takes as input the alignment and optional background knowledge and, using a voting algorithm between different linguistic-based heuristics, it aims at refining “ $\equiv$ ” semantic relations to “ $\sqsubseteq$ ”/“ $\sqsupseteq$ ” ones (and never the other way round).

(Weighted) Feedback Edge Set Problem. To the best of our knowledge, no suitable heuristics have been proposed for approximating the *WFES* problem. An heuristic based on *Simulated Annealing* has been proposed by Galinier *et al.* [22] for solving the *FVS* and *FES* problems. Unfortunately their approach cannot be trivially extended to *WFES*, due to their local search mechanism. On the contrary, exact and optimal approximations to the problem have been proposed [18].

Assumption of Disjointness. The assumption of disjointness has been originally introduced by Schlobach [67] to enhance the repair of ontologies that were underspecified in terms of disjointness axioms. In [56], a similar assumption is followed in the context of ontology mappings repair, where the authors restricted the number of disjointness axioms by using learning techniques [21, 82]. These techniques, however, typically require a manually created training set. In [20] the authors present an interactive system to guide the users in the manual enrichment of the ontologies with negative constraints, including disjointness axioms.

Our proposal, as [20, 21, 56, 82], aims at adding a small set of disjointness axioms, since adding all possible disjointness may be unfeasible for large ontologies. However, our method does not require manual intervention. Furthermore, to address the scalability problem when dealing with large ontologies and mapping sets, our method relies on the propositional projection of the input ontologies.

9. Conclusions

In this paper, we have presented a fully-automatic multi-strategy method to detect and correct conservativity violations in practice. We have extended the detection and repair algorithm for subsumption violations [75] with a repair algorithm expressed by means of an *ASP* program, tailored for equivalence violations. The conducted evaluation, significantly extending that of [74], supports the practical effectiveness of our approximate methods.

Our method is sound (the violations that are detected are indeed violations if considering the full expressiveness of the input ontologies), but incomplete, since the used approximate projections of the input ontologies (*i.e.*, Horn propositional and graph encodings) may lead to some violations being missed. In order to mitigate incompleteness, we plan to study extensions of our techniques to more expressive logical fragments, while keeping the current scalability properties. We also aim at adapting our techniques to cope with scenarios involving more than two input ontologies where new (repair) challenges arise [15].

The proposed algorithms follow a “better safe than sorry” approach, suitable for scenarios where the input ontologies are not modifiable. Nevertheless we plan to explore alternative methods to address the conservativity violations. For example, domain experts could be involved in the assessment of the additional disjointness [20, 35], and to suggest extensions to the input ontologies [31] for violations recognised as false positives.

Our techniques have been already deployed in Siemens [45] and Statoil [44] as part of BOOTOX [38], one of the components the “Ontology and mapping management module” provided by the Optique’s platform [24]. We consider, however, that the proposed methods have also potential in scenarios others than Optique. For instance, our approach has also been successfully employed in the context of multi-agent system (MAS) [40, 41] for assessing and favouring the compatibility among agents equipped with OWL 2 ontologies. The approach presented in [40, 41] is based on the novel inquiry dialogue (*i.e.*, a dialogue between two entities with the aim of solving a goal through knowledge exchange) introduced in [61, 62]. The dialogue aims at reusing existing mappings related to a domain of interest, with a minimal disclosure of private knowledge.²⁴ In this context, on the one hand, the ontology of the other agent(s) cannot be altered, on the other hand, modifying the local ontology after any dialogue could pose some risks. For example, consider the effect of a malicious agent forging ad-hoc mappings with the aim of altering the local ontology of the interlocutor. Different policies could be conceived in the dialogue. The family of approaches from Lambrix *et al.* [48, 50] could be referred to as *optimistic*, while forbidding alterations to the local ontology as *sceptical*. Of course mixed approaches could be conceived as well. A sceptical approach, as the one followed in [40, 41], could further exploit the techniques to detect violations of the consistency and conservativity principles in order to also estimate, under a game-theoretic point of view, the convenience of dealing with a particular agent, given the risk represented by these violations. In case of too many violations, the agent could simply decide to refuse to communicate, and seek for another (hopefully more compatible) agent.

Another approach, described in [53], also applies ontology matching in a MAS scenario in order to allow the exchange and extension of ontology-based *action plans* among agents. In such a context, violations of the conservativity principle should be taken into account and highly critical tasks should not be performed if such violations are detected.

Acknowledgements. Ernesto Jiménez-Ruiz was funded by the European Commission under FP7 Grant Agreement 318338, “Optique”, and the EPSRC projects Score!, ED3 and DBOnto. We also thank the unvaluable help provided by Bernardo Cuenca and Ian Horrocks. We are also very grateful for the support of the Optique colleagues that facilitated our understanding of the domain, especially: Dag Hovland, Evgeny Kharlamov, Dmitry Zheleznyakov, Martin Giese and Martin G. Skjæveland. Finally, we would also like to thank the anonymous reviewers of this paper.

References

- [1] Rakesh Agrawal, Alexander Borgida, and Hosagrahar V. Jagadish. Efficient Management of Transitive Relationships in Large Data and Knowledge Bases. In *ACM SIGMOD Conference on Management of Data*, pages 253–262, 1989.
- [2] Patrick Arnold and Erhard Rahm. Semantic Enrichment of Ontology Mappings: A Linguistic-Based Approach. In *Advances in Databases and Information System - East European Conference (ADBIS)*, pages 42–55, 2013.
- [3] F. Baader, I. Horrocks, and U. Sattler. Chapter 3 description logics. In Bruce Porter Frank van Harmelen, Vladimir Lifschitz, editor, *Handbook of Knowledge Representation*, volume 3 of *Foundations of Artificial Intelligence*, pages 135 – 179. Elsevier, 2008.
- [4] Elena Beisswanger and Udo Hahn. Towards Valid and Reusable Reference Alignments: Ten Basic Quality Checks for Ontology Alignments and their Application to Three Different Reference Data Sets. *Journal of Biomedical Semantics*, 3(Suppl 1):S4, 2012.

²⁴ Each agent is equipped with a local (private) OWL 2 ontology.

- [5] Olivier Bodenreider. The Unified Medical Language System (UMLS): integrating biomedical terminology. *Nucleic Acids Research*, 32:267–270, 2004.
- [6] Alexander Borgida and Luciano Serafini. Distributed Description Logics: Assimilating Information from Peer Sources. *Journal on Data Semantics*, 1:153–184, 2003.
- [7] Catia Pesquita and Daniel Faria and Emanuel Santos and Francisco M. Couto. To repair or not to repair: reconciling correctness and coherence in ontology reference alignments. In *Ontology Matching Workshop (OM)*, pages 13–24, 2013.
- [8] Vassilis Christophides, Dimitris Plexousakis, Michel Scholl, and Sotirios Tourounis. On Labeling Schemes for the Semantic Web. In *International World Wide Web Conference (WWW)*, pages 544–555, 2003.
- [9] Bernardo Cuenca Grau, Ian Horrocks, Yevgeny Kazakov, and Ulrike Sattler. Just the Right Amount: Extracting Modules from Ontologies. In *International Conference on World Wide Web (WWW)*, pages 717–726. ACM, 2007.
- [10] Bernardo Cuenca Grau, Ian Horrocks, Yevgeny Kazakov, and Ulrike Sattler. Modular Reuse of Ontologies: Theory and Practice. *Journal of Artificial Intelligence Research (JAIR)*, 31:273–318, 2008.
- [11] Bernardo Cuenca Grau, Ian Horrocks, Boris Motik, Bijan Parsia, Peter F. Patel-Schneider, and Ulrike Sattler. OWL 2: The next step for OWL. *Journal of Web Semantics*, 6(4):309–322, 2008.
- [12] Jérôme David, Jérôme Euzenat, François Scharffe, and Cássia Trojahn. The Alignment API 4.0. *Semantic Web Journal*, 2(1):3–10, 2011.
- [13] Agostino Dovier, Andrea Formisano, and Enrico Pontelli. An experimental comparison of constraint logic programming and answer set programming. In *AAAI*, volume 7, pages 1622–1625, 2007.
- [14] William F. Dowling and Jean H. Gallier. Linear-Time Algorithms for Testing the Satisfiability of Propositional Horn Formulae. *Journal of Logic Programming*, 1(3):267–284, 1984.
- [15] Jérôme Euzenat. Revision in networks of ontologies. *Artif. Intell.*, 228:195–216, 2015.
- [16] Jérôme Euzenat, Christian Meilicke, Heiner Stuckenschmidt, Pavel Shvaiko, and Cássia Trojahn. Ontology Alignment Evaluation Initiative: Six Years of Experience. *Journal on Data Semantics*, 15:158–192, 2011.
- [17] Jérôme Euzenat and Pavel Shvaiko. *Ontology Matching*. Springer, 2010.
- [18] Guy Even, J. Seffi Naor, Baruch Schieber, and Madhu Sudan. Approximating Minimum Feedback Sets and Multicuts in Directed Graphs. *Algorithmica*, 20(2):151–174, 1998.
- [19] Daniel Faria, Catia Pesquita, Emanuel Santos, Matteo Palmonari, Isabel F. Cruz, and Francisco M. Couto. The AgreementMakerLight Ontology Matching System. In *OTM Conferences*, pages 527–541, 2013.
- [20] Sébastien Ferré and Sebastian Rudolph. Advocatus Diaboli - Exploratory Enrichment of Ontologies with Negative Constraints. In *International Conference on Knowledge Engineering (EKAW)*, pages 42–56, 2012.
- [21] Daniel Fleischhacker and Johanna Völker. Inductive Learning of Disjointness Axioms. In *OTM Conferences*, pages 680–697, 2011.
- [22] Philippe Galinier, Eunice Lemamou, and Mohamed Wassim Bouzidi. Applying Local Search to the Feedback Vertex Set Problem. *Journal of Heuristics*, pages 1–22, 2013.
- [23] Giorgio Gallo and Giampaolo Urbani. Algorithms for Testing the Satisfiability of Propositional Formulae. *Journal of Logic Programming*, 7(1):45–61, 1989.
- [24] Martin Giese, Ahmet Soylu, Guillermo Vega-Gorgojo, Arild Waaler, Peter Haase, Ernesto Jiménez-Ruiz, Davide Lanti, Martín Rezk, Guohui Xiao, Özgür L. Özçep, and Riccardo Rosati. Optique: Zooming in on Big Data. *IEEE Computer*, 48(3):60–67, 2015.
- [25] Birte Glimm, Ian Horrocks, Boris Motik, Giorgos Stoilos, and Zhe Wang. Hermit: An OWL 2 reasoner. *J. Autom. Reasoning*, 53(3):245–269, 2014.
- [26] Jennifer Golbeck, Gilberto Fragoso, Frank W. Hartel, James A. Hendler, Jim Oberthaler, and Bijan Parsia. The National Cancer Institute’s Thesaurus and Ontology. *J. Web Sem.*, 1(1):75–80, 2003.
- [27] Rafael S Gonçalves, Bijan Parsia, and Ulrike Sattler. Concept-based semantic difference in expressive description logics. In *International Semantic Web Conference (ISWC)*, pages 99–115. Springer, 2012.
- [28] Matthew Horridge. *Justification Based Explanation in Ontologies*. PhD thesis, University of Manchester, 2011.
- [29] Matthew Horridge, Bijan Parsia, and Ulrike Sattler. Laconic and Precise Justifications in OWL. *International Semantic Web Conference (ISWC)*, pages 323–338, 2008.
- [30] Ian Horrocks, Oliver Kutz, and Ulrike Sattler. The even more Irresistible SROIQ. In *International Conference on Principles of Knowledge Representation and Reasoning (KR)*, pages 57–67, 2006.
- [31] Valentina Ivanova and Patrick Lambrix. A Unified Approach for Aligning Taxonomies and Debugging Taxonomies and their Alignments. In *European Semantic Web Conference (ESWC)*, pages 1–15. Springer, 2013.
- [32] Yves R. Jean-Mary, E. Patrick Shironoshita, and Mansur R. Kabuka. Ontology Matching With Semantic Verification. *Journal of Web Semantics*, 7(3):235–251, 2009.

- [33]Ernesto Jiménez-Ruiz and Bernardo Cuenca Grau. LogMap: Logic-based and Scalable Ontology Matching. In *International Semantic Web Conference (ISWC)*, pages 273–288, 2011.
- [34]Ernesto Jiménez-Ruiz, Bernardo Cuenca Grau, and Ian Horrocks. On the Feasibility of Using OWL 2 DL Reasoners for Ontology Matching Problems. In *OWL Reasoner Evaluation Workshop (ORE)*, 2012.
- [35]Ernesto Jiménez-Ruiz, Bernardo Cuenca Grau, Ian Horrocks, and Rafael Berlanga. Ontology integration using mappings: Towards getting the right logical consequences. In *European Semantic Web Conference (ESWC)*, pages 173–187, 2009.
- [36]Ernesto Jiménez-Ruiz, Bernardo Cuenca Grau, Ian Horrocks, and Rafael Berlanga. Logic-based Assessment of the Compatibility of UMLS Ontology Sources. *Journal of Biomedical Semantics*, 2(Suppl 1):S2, 2011.
- [37]Ernesto Jiménez-Ruiz, Bernardo Cuenca Grau, Yujiao Zhou, and Ian Horrocks. Large-scale Interactive Ontology Matching: Algorithms and Implementation. In *European Conference on Artificial Intelligence (ECAI)*, pages 444–449, 2012.
- [38]Ernesto Jiménez-Ruiz, Evgeny Kharlamov, Dmitriy Zheleznyakov, Ian Horrocks, Christoph Pinkel, Martin G. Skjæveland, Evgenij Thorstensen, and José Mora. BootOX: Practical Mapping of RDBs to OWL 2. In *International Semantic Web Conference (ISWC)*, pages 113–132, 2015.
- [39]Ernesto Jiménez-Ruiz, Christian Meilicke, Bernardo Cuenca Grau, and Ian Horrocks. Evaluating Mapping Repair Systems with Large Biomedical Ontologies. In *Description Logics (DL)*, pages 246–257, 2013.
- [40]Ernesto Jiménez-Ruiz, Terry R. Payne, Alessandro Solimando, and Valentina Tamma. Avoiding Alignment-based Conservativity Violations through Dialogue. In *International Workshop on OWL: Experiences and Directions (OWLED)*, 2015.
- [41]Ernesto Jiménez-Ruiz, Terry R. Payne, Alessandro Solimando, and Valentina A. M. Tamma. Limiting Logical Violations in Ontology Alignment Through Negotiation. In *Principles of Knowledge Representation and Reasoning: Proceedings of the Fifteenth International Conference (KR)*, pages 217–226, 2016.
- [42]Aditya Kalyanpur, Bijan Parsia, Matthew Horridge, and Evren Sirin. Finding all justifications of OWL DL entailments. *International Semantic Web Conference (ISWC)*, pages 267–280, 2007.
- [43]Yevgeny Kazakov, Markus Krötzsch, and Frantisek Simancik. The incredible ELK - from polynomial procedures to efficient reasoning with $\mathcal{EL}$ ontologies. *J. Autom. Reasoning*, 53(1):1–61, 2014.
- [44]Evgeny Kharlamov, Dag Hovland, Ernesto Jiménez-Ruiz, Davide Lanti, Hallstein Lie, Christoph Pinkel, Martin Rezk, Martin G. Skjæveland, Evgenij Thorstensen, Guohui Xiao, Dmitriy Zheleznyakov, and Ian Horrocks. Ontology Based Access to Exploration Data at Statoil. In *International Semantic Web Conference (ISWC)*, pages 93–112, 2015.
- [45]Evgeny Kharlamov, Nina Solomakhina, Özgür Lütfü Özçep, Dmitriy Zheleznyakov, Thomas Hubauer, Steffen Lamparter, Mikhail Roshchin, Ahmet Soylu, and Stuart Watson. How Semantic Technologies Can Enhance Data Access at Siemens Energy. In *International Semantic Web Conference (ISWC)*, pages 601–619, 2014.
- [46]Boris Konev, Dirk Walther, and Frank Wolter. The Logical Difference Problem for Description Logic Terminologies. In *International Joint Conference on Automated Reasoning (IJCAR)*, pages 259–274, 2008.
- [47]Roman Kontchakov, Frank Wolter, and Michael Zakharyashev. Can you Tell the Difference Between DL-Lite Ontologies? In *International Conference on Principles of Knowledge Representation and Reasoning (KR)*, 2008.
- [48]Patrick Lambrix, Zlatan Dragisic, and Valentina Ivanova. Get My Pizza Right: Repairing Missing Is-a Relations in $\mathcal{ALC}$ Ontologies. In *Semantic Technology*, pages 17–32. Springer, 2013.
- [49]Patrick Lambrix and Qiang Liu. Debugging the Missing Is-a Structure Within Taxonomies Networked by Partial Reference Alignments. *Data Knowledge Engineering (DKE)*, 86:179–205, 2013.
- [50]Patrick Lambrix, Fang Wei-Kleiner, Zlatan Dragisic, and Valentina Ivanova. Repairing Missing Is-a Structure in Ontologies is an Abductive Reasoning Problem. In *International Workshop on Debugging Ontologies and Ontology Mappings (WoDOOM)*, page 33, 2013.
- [51]Carsten Lutz, Dirk Walther, and Frank Wolter. Conservative Extensions in Expressive Description Logics. In *International Joint Conference on Artificial Intelligence (IJCAI)*, volume 7, pages 453–458, 2007.
- [52]Carsten Lutz and Frank Wolter. Deciding Inseparability and Conservative Extensions in the Description Logic EL. *Journal of Symbolic Computing*, 45(2):194–228, 2010.
- [53]Viviana Mascardi, Davide Ancona, Matteo Barbieri, Rafael H. Bordini, and Alessandro Ricci. Cool-AgentSpeak: Endowing AgentSpeak-DL Agents with Plan Exchange and Ontology Services. *Web Intelligence and Agent Systems*, 12(1):83–107, 2014.
- [54]Christian Meilicke. *Alignments Incoherency in Ontology Matching*. PhD thesis, University of Mannheim, 2011.
- [55]Christian Meilicke, Heiner Stuckenschmidt, and Andrei Tamin. Reasoning Support for Mapping Revision. *Journal of Logical Computing*, 19(5), 2009.

- [56]Christian Meilicke, Johanna Völker, and Heiner Stuckenschmidt. Learning Disjointness for Debugging Mappings between Lightweight Ontologies. In *International Conference on Knowledge Engineering (EKAW)*, pages 93–108, 2008.
- [57]Sergey Melnik, Hector Garcia-Molina, and Erhard Rahm. Similarity Flooding: A Versatile Graph Matching Algorithm and Its Application to Schema Matching. In *IEEE International Conference on Data Engineering (ICDE)*, pages 117–128, 2002.
- [58]Boris Motik, Bernardo Cuenca Grau, Ian Horrocks, and Ulrike Sattler. Representing Ontologies using Description Logics, Description Graphs, and Rules. *Artificial Intelligence Journal*, 173(14):1275–1309, 2009.
- [59]Victoria Nebot and Rafael Berlanga. Efficient Retrieval of Ontology Fragments Using an Interval Labeling Scheme. *Information Science Journal*, 179(24):4151–4173, 2009.
- [60]DuyHoa Ngo and Zohra Bellahsene. YAM++ results for OAEI 2013. In *Ontology Matching Workshop (OM)*, pages 211–218, 2013.
- [61]Terry R. Payne and Valentina Tamma. A Dialectical Approach to Selectively Reusing Ontological Correspondences. In *Knowledge Engineering and Knowledge Management (EKAW)*, pages 397–412. Springer, 2014.
- [62]Terry R. Payne and Valentina Tamma. Negotiating over Ontological Correspondences with Asymmetric and Incomplete Knowledge. In *Int'l Conference on Autonomous Agents and Multi-Agent Systems (AAMAS)*, pages 517–524, 2014.
- [63]Raymond Reiter. A Theory of Diagnosis from First Principles. *Artificial Intelligence Journal*, 32(1), 1987.
- [64]Mariano Rodriguez-Muro and Martín Rezk. Efficient SPARQL-to-SQL with R2RML mappings. *J. Web Sem.*, 33:141–169, 2015.
- [65]Cornelius Rosse and José L. V. Mejino Jr. A reference ontology for biomedical informatics: the Foundational Model of Anatomy. *J. Biomed. Informatics*, 36(6):478–500, 2003.
- [66]Emanuel Santos, Daniel Faria, Cátia Pesquita, and Francisco Couto. Ontology Alignment Repair Through Modularization and Confidence-based Heuristics. *arXiv:1307.5322 preprint*, 2013.
- [67]Stefan Schlobach. Debugging and Semantic Clarification by Pinpointing. In *European Semantic Web Conference (ESWC)*, pages 226–240. Springer, 2005.
- [68]Stefan Schlobach and Ronald Cornet. Non-standard Reasoning Services for the Debugging of Description Logic Terminologies. In *International Joint Conference on Artificial Intelligence (IJCAI)*, pages 355–362, 2003.
- [69]Stefan Schulz, Ronald Cornet, and Kent A. Spackman. Consolidating SNOMED CT’s ontological commitment. *Applied Ontology*, 6(1):1–11, 2011.
- [70]M. Berkan Sesen, René Banares-Alcántara, John Fox, Timor Kadir, and J. Michael Brady. Lung Cancer Assistant: an Ontology-Driven, Online Decision Support Prototype. In *International Workshop on OWL: Experiences and Directions (OWLED)*, 2012.
- [71]Pavel Shvaiko and Jérôme Euzenat. Ontology Matching: State of the Art and Future Challenges. *Transactions on Knowledge and Data Engineering (TKDE)*, 2012.
- [72]Martin G. Skjæveland, Espen H. Lian, and Ian Horrocks. Publishing the Norwegian Petroleum Directorate’s FactPages as Semantic Web Data. In *International Semantic Web Conference (ISWC)*, pages 162–177, 2013.
- [73]Alessandro Solimando. *Change Management in the Traditional and Semantic Web*. PhD thesis, University of Genoa, 2015. <https://github.com/asolimando/logmap-conservativity/raw/master/SolimandoA-thesis.pdf>.
- [74]Alessandro Solimando, Ernesto Jiménez-Ruiz, and Giovanna Guerrini. A Multi-strategy Approach for Detecting and Correcting Conservativity Principle Violations in Ontology Alignments. In *International Workshop on OWL: Experiences and Directions (OWLED)*, pages 13–24, 2014.
- [75]Alessandro Solimando, Ernesto Jiménez-Ruiz, and Giovanna Guerrini. Detecting and Correcting Conservativity Principle Violations in Ontology-to-Ontology Mappings. In *International Semantic Web Conference (ISWC)*, pages 1–16, 2014.
- [76]Alessandro Solimando, Ernesto Jiménez-Ruiz, and Giovanna Guerrini. On the Feasibility of Using OWL 2 Reasoners in Ontology Alignment Repair Problems. In *OWL Reasoner Evaluation Workshop (ORE)*, pages 60–67, 2015.
- [77]Alessandro Solimando, Ernesto Jiménez-Ruiz, and Giovanna Guerrini. Pushing the limits of OWL 2 reasoners in ontology alignment repair problems. *Intelligenza Artificiale*, 10:1–18, 2016.
- [78]Alessandro Solimando, Ernesto Jiménez-Ruiz, and Christoph Pinkel. Evaluating Ontology Alignment Systems in Query Answering Tasks. In *International Semantic Web Posters & Demonstrations Track (ISWC)*, pages 301–304, 2014.
- [79]Ahmet Soyulu, Martin Giese, Ernesto Jimenez-Ruiz, Guillermo Vega-Gorgojo, and Ian Horrocks. Experi-

- encing OptiqueVQS: A Multi-Paradigm and Ontology-Based Visual Query System for End Users. *Univ. Access in the Inform. Society*, 2015.
- [80]Boontawee Suntisrivaraporn, Guilin Qi, Qiu Ji, and Peter Haase. A Modularization-Based Approach to Finding All Justifications for OWL DL Entailments. In *Asian Semantic Web Conference (ASWC)*, pages 1–15, 2008.
- [81]Robert Tarjan. Depth-first Search and Linear Graph Algorithms. *SIAM Journal on Computing*, 1(2):146–160, 1972.
- [82]Johanna Völker, Denny Vrandečić, York Sure, and Andreas Hotho. Learning Disjointness. In *European Semantic Web Conference (ESWC)*, pages 175–189, 2007.
- [83]Ondřej Šváb, Vojtěch Svátek, Petr Berka, Dušan Rak, and Petr Tomášek. OntoFarm: Towards an Experimental Collection of Parallel Ontologies. In *International Semantic Web Conference (ISWC). Poster Session*, 2005.
- [84]W3C as Hitzler, Pascal and Krötzsch, Markus and Parsia, Bijan Patel-Schneider, Peter F. and Rudolph, Sebastian. OWL 2 Web Ontology Language Primer. <http://www.w3.org/TR/owl2-primer/>, 2009.
- [85]Peng Wang and Baowen Xu. Debugging Ontology Mappings: A Static Approach. *Computing and Informatics*, 27(1):21–36, 2012.
- [86]Songmao Zhang, Peter Mork, and Olivier Bodenreider. Lessons Learned from Aligning two Representations of Anatomy. In *International Conference on Principles of Knowledge Representation and Reasoning (KR)*, 2004.
- [87]Antoine Zimmermann and Jérôme Euzenat. Three Semantics for Distributed Systems and their Relations with Alignment Composition. In *International Semantic Web Conference (ISWC)*, pages 16–29. Springer, 2006.

A. Appendix

This section investigates the computational complexity of diagnosis computation (Section A.1), and its decomposability into subproblems (Section A.2).

A.1. Diagnosis Computation Complexity

With the aim of proving that *MAP-WFES* is *NP*-hard, Proposition A.1 introduces a polynomial reduction from *WFES* to *MAP-WFES*, denoted as $WFES \preceq MAP-WFES$. The intuitive idea behind the reduction is the following. Each arc (t, v, c) of the original graph is “split” in two arcs (t, u, c) and (u, v, c) , with u a fresh node. All the nodes t, v are associated with one of the input ontology, while the fresh nodes are associated with the other. In this way, all the arcs are mappings (*i.e.*, all of them are potentially removable, exactly as the original arcs). It is easy to see that the reduction preserves the solution weight and that a 1-1 correspondence exists between cycles in the two graphs. In addition, we remark that *MAP-WFES* does not break cycles traversing only vertices of one of the input ontologies. No such cycles can exist because all the arcs are mappings, as discussed above. A reduction example is given in Example A.1, followed by the definition of the reduction in Proposition A.1.

Example A.1. In Figure 13 graphs G (left) and G' (right) are shown. $WFES \preceq MAP-WFES$ coincides with G' . The solution to G' is equal to $\Delta = \{(c, cd, 1), (b, bg, 0.1), (gf, f, 0.4), (a, af, 0.2)\}$, with a total weight of 1.7. $D = \{(b, g, 0.1), (g, f, 0.4), (a, f, 0.2), (c, d, 1)\}$ is the corresponding solution to the instance of the *WFES* problem represented by G , and can be easily verified that is both minimal (having weight 1.7) and correct.

Proposition A.1. $WFES \preceq MAP-WFES$. A polynomial reduction from the *WFES* problem to the *MAP-WFES* problem exists. Let $G = (V, A)$ be a digraph. The reduction consists in constructing a digraph $G' = (V', A')$ such that a subset of edges, namely $\Delta \subseteq A$, is a solution to *MAP-WFES* iff the corresponding set of arcs, namely D , is a solution to *WFES* on G . The reduction is as follows:

1. for each $(x, y, c) \in A$, we create a fresh vertex v_{xy} , we add it to V'_2 , and we create a pair of arcs $(x, v_{xy}, c), (v_{xy}, y, c)$ that are added to A' and $\mathcal{M}$,
2. $V'_1 = V$ and $V' = V'_1 \cup V'_2$.

A set of arcs, namely $\Delta \subseteq A'$, is a solution to the *MAP-WFES* problem on digraph G' iff the corresponding feedback edge set D is a solution to the *WFES* G , where G' is computed from G , and for each arc of the form (x, v_{xy}, c) or (v_{xy}, y, c) in Δ , we have a corresponding arc (x, y, c) in D .

Proof. In order to prove the correctness of the reduction, we need to show that, if $G \preceq G'$, with G an instance of the *WFES* problem and G' an instance of the *MAP-WFES* problem, a set of arcs Δ is a (minimal) solution to G' iff the corresponding set of arcs D is a (minimal) solution to G . As discussed in [18], the proposed reduction is polynomial and it preserves graph connectivity and the weight of the solutions, by preserving a 1-1 correspondence between cycles of G and those of G' , due to the 1-1 correspondency between the arcs of A and those of A' .

$\Rightarrow$: If D is a solution to *WFES* on G , Δ is a solution to *MAP-WFES* on G' . Suppose that Δ is not a solution. This requires that, either at least a cycle κ' exists in digraph $(V', A' \setminus \Delta)$ or that a diagnosis Δ' exists such that $w(\Delta') < w(\Delta)$. For the

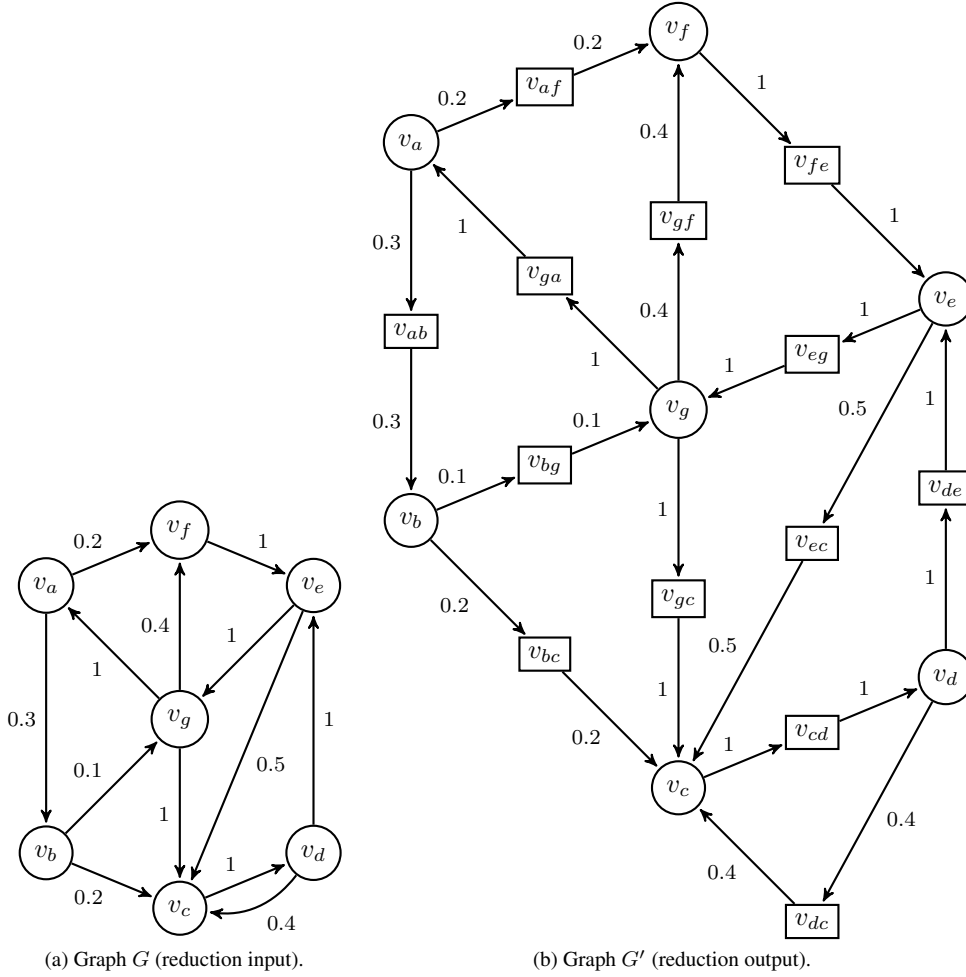


Fig. 13. Input graph G for $WFES$ (left) reduced to a corresponding graph G' , input for $MAP-WFES$ (right). In graph G' , all the round-shaped vertices belong to V'_1 , while square-shaped vertices belong to V'_2 .

first case, given the 1-1 correspondence between cycles of G and G' , this implies that a corresponding cycle in G exists as well, thus contradicting that D is a solution to the instance of the $WFES$ problem represented by G . For the latter case, given the 1-1 correspondence between arcs of G and G' , this implies that a solution D' corresponding to Δ' exists. By the weight preservation property of the reduction, $w(D') < w(D)$ holds, contradicting that D is a (minimal) solution to the instance of the $WFES$ problem represented by G .

$\Leftarrow$: If Δ is a solution to $MAP-WFES$ on G' , D is a solution to $WFES$ on G . Suppose that D is not. This requires that, either a cycle κ of G exists in digraph $(V, A \setminus D)$, or that a solution D' exists such that $w(D') < w(D)$. The first case requires the existence of a cycle κ' of G' (corresponding to κ) that is left unbroken. In turn, this either violates

that Δ is a solution, or that, for some $i \in \{1, 2\}$, κ' exclusively traverses elements of V'_i . This situation is excluded by construction of G' , because no arcs between vertices of the same subset V'_i of V' exist. For the latter case, this implies that a diagnosis Δ' corresponding to the (minimal) solution D' does not exist. This requires that at least an element of D' cannot belong to a diagnosis (*i.e.*, it cannot be removed). By construction of G' , we have that $\mathcal{M} = A'$. Given that only elements of $A' \setminus \mathcal{M}$ cannot belong to a diagnosis for the *MAP-WFES* problem, this results in a contradiction, thus proving the correctness of the reduction. $\square$

From the results of Proposition A.1 follows that *MAP-WFES* is *NP*-hard, as detailed in Proposition A.2.

Proposition A.2. *MAP-WFES* is *NP*-hard.

Proof. The proof follows from the polynomial reduction from the *WFES* problem, that is *NP*-hard [18], to *MAP-WFES*. $\square$

A.2. Decomposability of Equivalence Violations Diagnosis Computation

Proposition A.3 relates unsafe cycles and problematic SCCs, showing that each unsafe cycle results in a problematic SCC.

Proposition A.3. A SCC is problematic iff it (totally) contains at least one unsafe cycle.

Proof. $\Rightarrow$: Consider a problematic SCC S , with projections Π_1 and Π_2 on the input ontologies. From problematic SCC definition (Definition 4.6), at least one of the projections of S , say Π_1 , is not a local SCC. We therefore also know that Π_1 is not a SCC, otherwise it would also be a local SCC. Suppose that Π_1 is a subset of a SCC Π' . This implies that all the elements of Π' belongs to S as well, and therefore Π' and Π_1 are identical, but this contradicts the assumption that Π_1 is not a SCC.

$\Leftarrow$: from the definition of cycle and SCC, each cycle κ is contained in a SCC S (*i.e.*, $\kappa \subseteq S$). Let κ be an unsafe cycle, let also κ_1 (resp. κ_2) be the subset of vertices of κ belonging to an input ontology $\mathcal{O}_1$ (resp. $\mathcal{O}_2$). By definition of unsafe cycle (Definition 4.3), at least one of this subsets, say κ_1 , is not contained in any local SCC. But given that $\kappa \subseteq S$, $\kappa_1 \subseteq \Pi_{\mathcal{O}_1}(S)$ holds. Therefore, $\Pi_{\mathcal{O}_1}(S)$ is not contained in any local SCC either and, by Definition 4.6, S is a problematic SCC. $\square$

Proposition A.3 guarantees completeness for a detection technique for violations to the conservativity principle on a graph representation of an aligned ontology, based on problematic SCCs. Given that all the violations result in unsafe cycles, and that they totally belong to a single problematic SCC, completeness for a repair technique breaking all the unsafe cycles follows.

Notice also that a (unsafe) cycle always belongs to one and only one (problematic) SCC (as expressed by Proposition A.4), while a problematic SCC may contain more than one cycle. Therefore, a technique detecting problematic SCCs may be more efficient than one directly addressing unsafe cycles.

Proposition A.4. Safe cycle never traverse multiple SCCs of the same input ontology.

Proof. By Definition 4.3, all the vertices belonging to a projection Π of a safe cycle κ^s needs to be traversed by a cycle κ' in the input ontology these vertices belongs to. The claim is that cycle κ' identifies either a SCC of the aligned ontology, or a subset of a SCC. Assume that vertices of Π belong to at least two SCCs S_1, S_2 , that is, $\Pi \cap S_1 \neq \emptyset$ and $\Pi \cap S_2 \neq \emptyset$. Being traversed by a cycle, all the vertices of Π are mutually reachable.

Then, from transitivity of reachability, it follows that all the vertices in $S_1 \cup S_2$ are mutually reachable. This contradicts the hypothesis that S_1 and S_2 are two distinct SCCs, thus proving the proposition. Such argument can be straightforwardly generalized to more than two SCCs. $\square$

Proposition A.5 proves the correctness of our approach and the optimality of the computed (global) diagnosis.

Proposition A.5. Computing a (global) diagnosis for a graph G , representing an aligned ontology, can be reduced to computing the (local) diagnoses for the problematic SCCs of G . The (minimal) global diagnosis is the union of the (minimal) local diagnoses.

Proof. From Proposition A.3 it follows that: (i) all and only problematic SCCs contain unsafe cycles, (ii) an unsafe cycle does not traverse vertices of more than one SCC (i.e., the unsafe cycles of distinct SCCs are totally disjoint). From (i) completeness follows (it is sufficient to compute a diagnosis for each problematic SCC to remove all the unsafe cycles in the aligned ontology). (ii) ensures the independence of SCCs, and this guarantees minimality and correctness for local diagnoses computed in isolation. Finally, it is immediate to see that the minimality property is preserved by the union of local diagnoses, and this concludes the proposition. $\square$

Proposition A.5 thus guarantees that the global diagnosis computed as the union of the diagnoses of the problematic SCCs is both minimal and correct (that is, it breaks all the unsafe cycles).

Correspondence and offprint requests to: Alessandro Solimando,
Address: DIBRIS, University of Genova, Via Dodecaneso 35, 16146, Genova, Italy
Email: alessandro.solimando@unige.it