

Limiting Logical Violations in Ontology Alignment Through Negotiation

Ernesto Jiménez-Ruiz

Dept. Computer Science,
University of Oxford, UK

Terry R. Payne

Dept. Computer Science,
University of Liverpool, UK

Alessandro Solimando

DIBRIS,
Università di Genova, Italy

Valentina Tamma

Dept. Computer Science,
University of Liverpool, UK

Abstract

Ontology alignment (also called ontology matching) is the process of identifying *correspondences* between entities in different, possibly heterogeneous, ontologies. Traditional ontology alignment techniques rely on the full disclosure of the ontological models; however, within open and opportunistic environments, such approaches may not always be pragmatic or even acceptable (due to privacy concerns). Several studies have focussed on collaborative, decentralised approaches to ontology alignment, where agents negotiate the acceptability of single correspondences acquired from past encounters, or try to ascertain novel correspondences on the fly. However, such approaches can lead to logical violations that may undermine their utility. In this paper, we extend a dialogical approach to correspondence negotiation, whereby agents not only exchange details of possible correspondences, but also identify potential violations to the *consistency* and *conservativity* principles. We present a formal model of the dialogue, and show how agents can repair logical violations during the dialogue by invoking a correspondence repair, thus negotiating and exchanging *repair plans*. We illustrate this opportunistic alignment mechanism with an example and we empirically show that allowing agents to strategically reject or weaken correspondences when these cause violations does not degrade the effectiveness of the alignment computed, whilst reducing the number of residual violations.

Introduction

Autonomous agents rely on an internal representation, or *world model*, of their perceptions in order to behave appropriately in uncertain or unknown environments. This representation is often defined within some logical theory (*ontology*) that is completely or only partially shared with other agents, even though there may be common assumptions regarding how pertinent information and knowledge is modelled, expressed and interpreted. When interoperability between heterogeneous systems is required, an integration phase is necessary to reconcile different knowledge models and clarify implicit assumptions, especially within dynamic and opportunistic scenarios (*e.g.*, e-commerce, open-data or mobile systems).

Traditionally, the challenge of resolving semantic heterogeneity has been addressed by aligning the agents' on-

tologies, using one of many existing alignment systems (Shvaiko and Euzenat 2013; Cheatham and others 2015). However, most approaches are centralised, whereby one agent (or a third party) is responsible for generating alignments and has full access to the ontologies. Furthermore, no single approach is necessarily suitable for all scenarios; and (partial) privacy has become increasingly pertinent, whereby neither agent or knowledge system is prepared to disclose its full ontology (Cuenca Grau and Motik 2012; Payne and Tamma 2014b), *e.g.*, if the knowledge encoded within an ontology is confidential or commercially sensitive.

Recent alignment approaches have assumed the existence of pre-computed alignments (Laera et al. 2007; Trojahn dos Santos, Quaresma, and Vieira 2008; Doran et al. 2009; Meilicke 2011), from which new ones are selected and combined (*alignment aggregation*). These approaches exploit additional information, *e.g.*, the level of confidence or *weight* associated with each correspondence, estimated by its frequency. Furthermore, different alignment systems may map entities from one ontology to different entities in the other ontology, leading to *ambiguity*. Including such correspondences may be legitimate in certain scenarios; users may be familiar with the notion of synonyms or equivalent labels for certain concepts, and may not want to converge on a single canonical label. However, there is the danger that integrating such ambiguity within either ontology can lead to many undesirable logical consequences, and violate the three principles proposed by Jiménez-Ruiz et al. (2011): *consistency*, *locality*, and *conservativity*. In order to ensure injective alignments, many alignment systems employ a *brute-force* approach to the selection of unambiguous correspondences, through the identification of a *Matching* from the resulting weighted bipartite graph (obtained by mapping all the entities in the signatures of the two ontologies that are being aligned). This is done by finding either a maximum weighted bipartite matching which can be solved in polynomial time (Kuhn 1955), or finding a stable solution based on the *Stable Marriage* problem (Gale and Shapley 1962). Although these approaches are effective in eliminating ambiguous correspondences, they can also prune out potentially useful alternatives that satisfy the *conservativity principle*, where correspondences should not introduce new semantic relationships between (named) concepts from one of the input ontologies.

In this paper, we extend an existing decentralised approach (Payne and Tamma 2014b) to compute alignments, whereby agents engage in a dialogue to exchange details of possible correspondences, and identify and eliminate those potential ones that could yield conservativity and consistency violations by means of a detection and repair mechanism based on the approach presented in (Jiménez-Ruiz and Cuenca Grau 2011; Solimando, Jiménez-Ruiz, and Guerrini 2014a; 2014b). The dialogue represents a negotiation mechanism that allows the agents to rationally agree over a set of correspondences that 1) they jointly consider correct as they are above a given *admissibility threshold* and 2) do not cause any logical violation to either of the two agents' ontologies. This approach assumes that the agents had acquired correspondences from past encounters, or from publicly available alignment systems (that were kept private), and that each agent associated some *weight* to each known correspondence. As this knowledge is *asymmetric* and *incomplete* (i.e., neither agent involved in the dialogue is aware of all of the correspondences, and the weight assigned to each correspondence could vary greatly), the agents engage in the dialogue to: 1) ascertain the joint acceptability of each correspondence; and to 2) select a set of correspondences which reduce or eliminate the occurrence of possible conservativity and consistency violations (from each agent's individual perspective, rather than from a joint perspective). After introducing the main theory behind alignment repair, we present a formal model of the dialogue, show how conservativity and consistency violations can be repaired through the exchange of *repair plans*, and present a walkthrough example of a dialogue. We discuss termination and soundness of the dialogue and we empirically evaluate its effectiveness with respect to completeness. We finalise by summarising related work and drawing concluding remarks.

Preliminaries

In this section, we introduce the formal representation of ontology correspondences, and the notions of semantic difference, consistency principle violation, conservativity principle violation and alignment repair.

Representation of Ontology Correspondences

Agents can negotiate over the viability of different correspondences that could be used to align the two agents' ontologies. We assume that each agent commits to an *ontology* $\mathcal{O}$, which is an explicit and formally defined vocabulary representing the agent's knowledge about the environment, and its background knowledge (domain knowledge, beliefs, tasks, etc.). $\mathcal{O}$ is modeled as a set of axioms describing classes and the relations existing between them¹ and $\Sigma \subseteq \Sigma(\mathcal{O})$ is the public *ontology signature*; i.e., the set of class and property names used in $\mathcal{O}$ that is public (i.e. that an agent is prepared to disclose to other agents). To avoid confusion, the sender's ontology is denoted $\mathcal{O}^x$, whereas the recipient's ontology is $\mathcal{O}^{\hat{x}}$. For agents to interoperate in an encounter, they need to determine an *alignment* $\mathcal{A}$ between the two vocabulary fragments Σ^x and $\Sigma^{\hat{x}}$ for that encounter.

¹Here we restrict the ontology definition to classes and roles.

An alignment (Euzenat and Shvaiko 2013) consists of a set of *correspondences* that establish a logical relationship between the entities belonging to each of the two ontologies, and a set of logical relations. The universe of all possible correspondences is denoted $\mathcal{C}$. The aim of the dialogue is to generate an alignment $\mathcal{A} \subseteq \mathcal{C}$, that maps between the entities in Σ^x and $\Sigma^{\hat{x}}$, that does not introduce any conservativity or consistency violations, and whose joint weight is at least as great as an *admissibility threshold* ϵ that both agents assume.

Definition 1: A **correspondence** is a triple denoted $c = \langle e, e', r \rangle$ such that $e \in \Sigma^x$, $e' \in \Sigma^{\hat{x}}$, $r \in \{\equiv, \sqsubseteq, \sqsupseteq\}$.

Correspondences are usually formally represented as OWL 2 axioms to enable the reuse of the extensive range of available OWL 2 reasoning infrastructure, but alternative formal semantics for ontology correspondences have been proposed (for example, see Borgida and Serafini (2003)).

Semantic Consequences of the Integration

The ontology ($\mathcal{O}^{\mathcal{A}}$) resulting from the integration of two ontologies $\mathcal{O}^x$ and $\mathcal{O}^{\hat{x}}$ via a set of correspondences $\mathcal{A}$ may entail axioms that do not follow from $\mathcal{O}^x$, $\mathcal{O}^{\hat{x}}$, or $\mathcal{A}$ alone. These new semantic consequences can be captured by the notion of *deductive difference* (Konev, Walther, and Wolter 2008). Intuitively, the deductive difference between $\mathcal{O}$ and $\mathcal{O}'$ w.r.t. a signature Σ (i.e., set of entities) is the set of entailments constructed over Σ that do not hold in $\mathcal{O}$, but do hold in $\mathcal{O}'$. Unfortunately, no algorithm is available for computing deductive difference for DLs more expressive than $\mathcal{EL}$, for which the existence of tractable algorithms is still an open problem (Konev, Walther, and Wolter 2008). In order to avoid these limitations, practical applications typically rely on *approximations* of the deductive difference. For example, an approximation that only requires comparing the (atomic) classification hierarchies of $\mathcal{O}$ and $\mathcal{O}'$ provided by an OWL 2 reasoner is given in the following definition:

Definition 2: Let A, B be atomic concepts (including $\top, \perp$), Σ be a signature, $\mathcal{O}$ and $\mathcal{O}'$ be two OWL 2 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) $\mathcal{O}' \models a \sqsubseteq b$.

In this paper we rely on this approximation, which has successfully been used in the past in the context of ontology integration (Jiménez-Ruiz et al. 2009, inter alia).

Consistency and Conservativity Violations

The consistency principle requires that the vocabulary in $\mathcal{O}^{\mathcal{A}} = \mathcal{O}^x \cup \mathcal{O}^{\hat{x}} \cup \mathcal{A}$ be satisfiable, assuming the union of input ontologies $\mathcal{O}^x \cup \mathcal{O}^{\hat{x}}$ (without the alignment $\mathcal{A}$) does not contain unsatisfiable concepts.

Definition 3: An alignment $\mathcal{A}$ violates the **consistency principle** (i.e., it is *incoherent*) with respect to $\mathcal{O}^x$ and $\mathcal{O}^{\hat{x}}$, if $\text{diff}_{\Sigma}^{\approx}(\mathcal{O}^x \cup \mathcal{O}^{\hat{x}}, \mathcal{O}^{\mathcal{A}})$ contains axioms of the form $a \sqsubseteq \perp$, for any $a \in \Sigma = \Sigma(\mathcal{O}^x \cup \mathcal{O}^{\hat{x}})$. Violations of the consistency principle result in an *incoherent integrated ontology* $\mathcal{O}^{\mathcal{A}}$.

The consistency principle has been widely investigated in the literature, and approaches for detecting and repairing

logical inconsistencies in integrated ontologies have been proposed (Meilicke 2011; Jiménez-Ruiz et al. 2011).

The conservativity principle (general notion) states that the ontology $\mathcal{O}^A$ should not induce any change in the concept hierarchies of the input ontologies $\mathcal{O}^x$ and $\mathcal{O}^{\hat{x}}$. That is, the sets $\text{diff}_{\Sigma^x}^{\approx}(\mathcal{O}^x, \mathcal{O}^A)$ and $\text{diff}_{\Sigma^{\hat{x}}}^{\approx}(\mathcal{O}^{\hat{x}}, \mathcal{O}^A)$ must be empty for signatures Σ^x and $\Sigma^{\hat{x}}$, respectively. In this paper we use two less restrictive variants of this principle, which were presented by (Solimando, Jiménez-Ruiz, and Guerrini 2014a).

Definition 4: Let $\mathcal{O}$ be one of the input ontologies ($\mathcal{O}^x$ or $\mathcal{O}^{\hat{x}}$) and $\Sigma = \Sigma(\mathcal{O}) \setminus \{\top, \perp\}$ be its signature, let $\mathcal{A}$ be a coherent alignment between $\mathcal{O}^x$ and $\mathcal{O}^{\hat{x}}$, let $\mathcal{O}^A$ be the integrated ontology, and let a, b be atomic concepts in Σ . We define two sets of **conservativity violations** of $\mathcal{O}^A$ w.r.t. $\mathcal{O}$:²

- *subsumption violations*, denoted $\text{subViol}(\mathcal{O}, \mathcal{O}^A)$, as the set of $a \sqsubseteq b$ axioms satisfying: (i) $a \sqsubseteq b \in \text{diff}_{\Sigma}^{\approx}(\mathcal{O}, \mathcal{O}^A)$, (ii) $\mathcal{O} \not\models b \sqsubseteq a$, and (iii) there is no d in Σ s.t. $\mathcal{O} \models d \sqsubseteq a$, and $\mathcal{O} \models d \sqsubseteq b$ (i.e., no shared descendants for a and b).
- *equivalence violations*, denoted $\text{eqViol}(\mathcal{O}, \mathcal{O}^A)$, as the set of $a \equiv b$ axioms satisfying: (i) $\mathcal{O}^A \models a \equiv b$, (ii) $a \sqsubseteq b \in \text{diff}_{\Sigma}^{\approx}(\mathcal{O}, \mathcal{O}^A)$ and/or $b \sqsubseteq a \in \text{diff}_{\Sigma}^{\approx}(\mathcal{O}, \mathcal{O}^A)$.

Thus, an alignment $\mathcal{A}$ violates the **conservativity principle** if $\text{subViol}(\mathcal{O}, \mathcal{O}^A)$ or $\text{eqViol}(\mathcal{O}, \mathcal{O}^A)$ are not empty.

Note that the notion of subsumption violations relies on the assumption of disjointness (Schlobach 2005) and it can be reduced to a consistency principle problem.

Alignment Repair

An alignment $\mathcal{A}$ that violates the consistency and/or the conservativity principles can be fixed by removing correspondences from $\mathcal{A}$. This process is referred to as *alignment repair* (or repair for short).

A trivial repair is $\mathcal{R} = \mathcal{A}$, since an empty set of correspondences cannot lead to violations, according to Definitions 3 and 4. Nevertheless, the objective is to remove as few correspondences as possible. Minimal repairs can be computed by extracting the justifications for the violations (e.g., (Kalyanpur et al. 2007)), and selecting a hitting set of correspondences to be removed, following a minimality criteria (e.g., the number of removed correspondences). However, justification-based techniques do not scale when the number of violations is large (a typical scenario in alignment repair problems (Solimando, Jiménez-Ruiz, and Guerrini 2015)).

To address this scalability issue, alignment repair systems usually compute an *approximate repair* using incomplete reasoning techniques. Approximate repairs do not guarantee that $\mathcal{A} \setminus \mathcal{R}$ is violation free, but (in practice) it significantly minimizes the number of violations caused by the original alignment $\mathcal{A}$. In this paper, we rely on the (approximate) repair techniques presented in (Jiménez-Ruiz and Cuenca Grau 2011; Solimando, Jiménez-Ruiz, and Guerrini 2014a; 2014b) to minimize the violations to the consistency and conservativity principles. The detection and correction of consistency and (conservativity) subsumption violations (which are reduced to a consistency violations) is based on

the Dowling-Gallier algorithm (Dowling and Gallier 1984) for propositional Horn satisfiability; whereas equivalence violations are addressed using a combination of graph theory and logic programming. Both approaches 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 misreported. Nevertheless, incompleteness is mitigated thanks to the classification of the input ontologies using full reasoning.

The Dialogue

In (Payne and Tamma 2014b), the *Correspondence Inclusion Dialogue* was presented which enabled two agents to exchange knowledge about ontological correspondences resulting in an alignment that satisfies the following:

1. An agent is aware of a set of correspondences, each with an associated *weight* that represents the level of confidence the agent has in the correctness of the correspondence;
2. There should be no *ambiguity* w.r.t. either the source entities in the resulting alignment, or the target entities;
3. If there are alternative choices of correspondences, the selection is based on the combined, or *joint weight* of the correspondences (i.e., based on their combined weights ascribed by both agents);
4. No correspondences should be selected if their joint weight is less than an *admissibility-threshold*; and
5. the number of correspondences disclosed (i.e., whose weight is shared in the dialogue) should be *minimised*.

The rationale behind the dialogue exploited the fact that whilst the agents involved sought to minimise the disclosure of their ontological knowledge (and thus the concepts known), some exchange of ontological knowledge (at least the exchange of a subset of candidate correspondences) was necessary to determine a consensual set of correspondences that formed the final alignment. Whilst it was assumed that the agents were inherently self interested, there was also the assumption that the agents were collaborative with respect to determining an alignment that could facilitate communication (Grice 1975), as it was in the interest of all rational agents involved to be able to communicate successfully.

The dialogue has been significantly modified to retain the ability to negotiate over private weights in such a way as to minimise the number of correspondences disclosed, but to replace the ambiguity mechanism (i.e., removing the *object* move, and the use of argumentation) with a conservativity and consistency violation detection and repair mechanism (discussed in the next section). Furthermore, the formal treatment of correspondences and weights, as well as the syntax for the moves has been improved. The new dialogue is described in detail in the following subsections.

The Inquiry Dialogue Moves

The dialogue consists of a sequence of communicative acts, or *moves*, whereby agents take turns to assert the candidacy

²We assume that $\text{diff}_{\Sigma}^{\approx}(\mathcal{O}, \mathcal{O}^x \cup \mathcal{O}^{\hat{x}}) = \emptyset$

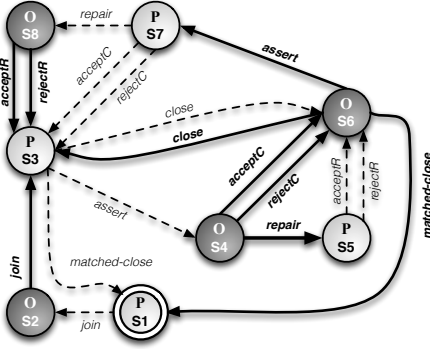


Figure 1: The dialogue as a state diagram. Nodes indicate the agent whose turn it is to utter a move. Moves uttered by the *proponent* P , i.e., the agent starting the dialogue, are labelled with a light font / dashed edge, whereas those uttered by *opponent* O are labelled with a heavy font / solid edge.

of some correspondence c for inclusion in a mutually acceptable alignment, $\mathcal{A}$. Agents associate a private, static *weight* κ_c to a correspondence c (where $0 \leq \kappa_c \leq 1$) that represents the confidence an agent has in the correctness of c . Each agent manages a private knowledge base, known as the *Correspondence Store* (Δ), which stores the correspondences and their associated private weights, and a public knowledge base, or *Commitment Store*, CS , which contains a trace of the moves uttered by both agents (Walton and Krabbe 1995). Although each agent maintains its own copy of the CS , these will always be identical, and thus we do not distinguish between them. We distinguish between the sender's Correspondence Store, Δ^x , and recipient's store, $\Delta^{\hat{x}}$.

In the dialogue, agents take turn in asserting correspondences and respond to such assertions by 1) confirming the acceptability of c without the need for any alignment repair; 2) proposing a possible repair to $\mathcal{A}$ to allow c to be added to $\mathcal{A}$ without introducing any consistency or conservativity violations (such as weakening or removing an existing correspondence); or 3) rejecting the acceptability of c . Each agent discloses its private *belief* regarding some correspondence c and its weight, and the agents negotiate to rationally identify a set of mutually acceptable correspondences, given an *admissibility threshold* ϵ . It assumes that only two agents (referred to here as *Alice* and *Bob* for clarity) participate in the dialogue, and that each agent plays a specific role (i.e., an agent is either a *sender* x or *recipient* $\hat{x}$) in any single dialogue move. As the dialogue is non deterministic and either *Bob* or *Alice* can start by issuing the first join move, in the state diagram illustrated in Figure 1 we refer to *proponent* and *opponent*, where the *proponent*, P , is the agent initiating the dialogue, and the *opponent*, O , is the agent that interacts with P in order to compute a final alignment. A state transition occurs when either P or O (whichever is currently playing the role of the sender x) sends a message to the recipient $\hat{x}$. Therefore, agents take turns in uttering *assert* moves (i.e., to transition from state S3 for P as the sender or S6 for O as the sender in the state diagram. A sender x can also make two consecutive moves in certain circumstances, such as af-

ter an *accept* or *reject* move (see states labelled S7 for P and S4 for O in Figure 1), to ensure they take turns in asserting new correspondences.

The set of possible *moves* $\mathcal{M}$ permitted by the dialogue are summarised in Table 1. The syntax of each move at time s is of the form $m_s = \langle x, \tau, c, \kappa_c, \mathcal{R} \rangle$, where x represents the identity of the agent making the move; $\tau \in \mathcal{M}$ represents the move type; c is the subject of the move, i.e., the correspondence that is being discussed; κ_c represents either the private or joint weight associated with c where $0 \leq \kappa_c \leq 1$; and $\mathcal{R}$ represents a repair for correspondences within the candidate alignment $\mathcal{A}$ or the correspondence c itself. For some moves, it may not be necessary to specify a correspondence, weight or repair; in which case they will be empty or unspecified (represented as *nil*).

Aggregating Weights and the Upper Bound

Within the dialogue, the agents try to ascertain the unambiguous, mutually acceptable correspondences to include in the final alignment $\mathcal{A}$ by selectively sharing those correspondences that are believed to have the highest weight. Once each agent knows of the other agent's weight for a correspondence c , it can calculate its *joint weight*, and check if it is greater than or equal to the *admissibility threshold*, ϵ . This threshold filters out correspondences with a low weight (i.e., when $\kappa_c < \epsilon$), whilst minimising the number of beliefs disclosed. The function $\text{joint} : \mathcal{C} \mapsto [0, 1]$ returns the *joint weight* for some correspondence $c \in \mathcal{C}$. This results in either: 1) κ_c^{joint} calculated based on the weights for both agents (if both weights are known); or 2) κ_c^{est} for a conservative upper estimate, if only one of the weights is known.

When the sender x receives an *assert* move from $\hat{x}$ for a correspondence it knows (i.e., where $c \in \Delta^x$), it can assess the joint weight for c as the average between its own weight and that shared by $\hat{x}$ (**Case 1**). If, however, x has no prior knowledge of c (i.e., $c \notin \Delta^x$), then the acceptability of the correspondence, and its joint weight will depend only on $\kappa_c^{\hat{x}}$ (**Case 2**). Finally, if x holds a belief on c that has not yet been disclosed to $\hat{x}$ ($c \in \Delta^x; c \notin CS$) and if $\kappa_c^{\hat{x}}$ has not been disclosed by $\hat{x}$, then an *upper bound* κ_u^x estimate is assumed (**Case 3**). The *upper bound*, κ_u^x is explained below.

Definition 5: The function $\text{joint} : \mathcal{C} \mapsto [0, 1]$ returns the joint weight for $c \in \mathcal{C}$:

$$\text{joint}(c) = \begin{cases} \text{avg}(\kappa_c^x, \kappa_c^{\hat{x}}) & \text{Case 1: } c \in \Delta^x \cap \Delta^{\hat{x}}, c \in CS \\ \kappa_c^{\hat{x}} & \text{Case 2: } c \notin \Delta^x, c \in CS \\ \text{avg}(\kappa_c^x, \kappa_u^x) & \text{Case 3: } c \in \Delta^x, c \notin CS \end{cases}$$

Each agent takes turns to propose a correspondence, and the other participant confirms if the joint weight $\kappa_c^{\text{joint}} \geq \epsilon$. Proposals are made by identifying an undisclosed correspondence with the highest weight κ_c^x . As the dialogue proceeds, each subsequent correspondence asserted will have an equivalent or lower weight than that previously asserted by the same agent.

Whenever a correspondence is asserted, the agent should check that its estimated joint weight κ_c^{est} is not less than the *admissibility threshold*, ϵ . Because the estimate is an upper estimate, the final joint weight κ_c^{joint} could subsequently be

Table 1: The set $\mathcal{M}$ of legal moves permitted by the dialogue.

Syntax	Description
$\langle x, \text{join}, \text{nil}, \text{nil}, \text{nil} \rangle$	Agents utter the <i>join</i> move to participate within a dialogue.
$\langle x, \text{assert}, c, \kappa_c^x, \mathcal{R} \rangle$	The agent x will <i>assert</i> the correspondence c that is believed to be viable for inclusion into the final alignment $\mathcal{A}$, and is the undisclosed correspondence with highest individual weight κ_c^x . If c violates conservativity or consistency given $\mathcal{A} \cup \mathcal{O}^x$, then $\mathcal{R}$ will contain an appropriate repair plan to solve this violation to be applied either to the correspondences already in $\mathcal{A}$ or to the newly asserted correspondence c .
$\langle x, \text{rejectC}, c, \text{nil}, \text{nil} \rangle$	If the c asserted in the previous move was not viable (i.e., $\kappa_c^{\text{joint}} < \epsilon$, or a violation was subsequently found for x where the repair involved removing c from $\mathcal{A}$), then it is rejected.
$\langle x, \text{acceptC}, c, \kappa_c^{\text{joint}}, \text{nil} \rangle$	Given c received in the previous <i>assert</i> move and the associated repair $\mathcal{R}$, if $\kappa_c^{\text{joint}} \geq \epsilon$ (i.e., the joint weight is above threshold), and no violation is generated for $\mathcal{A}' \cup \mathcal{O}^x$ (where $\mathcal{A}'$ is the result of applying the repairs in $\mathcal{R}$ to $\mathcal{A}$), then c is accepted and this joint weight is shared. If a further violation occurs due to $\mathcal{A}' \cup \mathcal{O}^x$, an additional repair should be generated and shared using the <i>repair</i> move.
$\langle x, \text{repair}, c, \kappa_c^{\text{joint}}, \mathcal{R} \rangle$	If the agent detected a violation in $\mathcal{A}' \cup \mathcal{O}^x$ (see <i>acceptC</i>), it utters a <i>repair</i> move to: 1) indicate that c is acceptable if $\hat{x}$ accepts the repair $\mathcal{R}$; and 2) inform $\hat{x}$ of the resulting joint weight κ_c^{joint} . Note that the previous repair will not be applied to $\mathcal{A}$ at this stage, as it is predicated on $\hat{x}$ accepting the repair.
$\langle x, \text{rejectR}, c, \text{nil}, \text{nil} \rangle$	If the agent rejects the proposed repair (e.g., it weakens or removes a mapping deemed necessary for its later transaction), then it can reject the repair, which will also result in c being implicitly rejected.
$\langle x, \text{acceptR}, c, \text{nil}, \text{nil} \rangle$	The agent x accepts the repair $\mathcal{R}$ for the correspondence c and updates $\mathcal{A}$. On receipt of this move, agent $\hat{x}$ also updates its version of $\mathcal{A}$.
$\langle x, \text{close}, \text{nil}, \text{nil}, \text{nil} \rangle$	An agent utters a <i>close</i> move when it has no more undisclosed viable candidate correspondences. The dialogue terminates when both agents utter a sequence of <i>close</i> moves (i.e., a <i>matched-close</i>).

lower, and the correspondence still rejected. Agents determine this upper estimate by exploiting the fact that assertions are always made on the undisclosed correspondence with the highest weight. Thus, if one agent asserts some correspondence, the other agent’s weight for that asserted correspondence will never be greater than their own previous assertion. Therefore, each agent maintains an *upper bound*, κ_u^x , corresponding to the other agents assertions (prior to the dialogue, $\kappa_u^x = 1.0$).

Strategic Generation of Repairs

The goal of each agent is to extend each ontology so that there exists a set of entities that are common to both ontologies (i.e., those in $\mathcal{A}$). This should subsequently facilitate the meaningful exchange of knowledge between the two agents, provided that it is expressed using entities within $\mathcal{A}$.

When a new correspondence is proposed during the dialogue, each agent verifies the suitability of this correspondence with respect to its commitment store and its admissibility threshold. The suitability with respect to the commitment store ensures that every addition to the set of correspondences already agreed in the alignment negotiated until that moment does not introduce any *conservativity* and *consistency* violation.

The violation detection and repair mechanism introduced earlier and defined in (Solimando, Jiménez-Ruiz, and Guerini 2014a; 2014b) has been adapted to *incrementally* check for consistency and conservativity violations as new correspondences are proposed for inclusion within $\mathcal{A}$.³ As the ontologies themselves are considered immutable, repairs can

³Note that we identify and repair consistency prior to conservativity violations, as unsatisfiable concepts would be subsumed by any other concept, thus leading to a very large number of (misleading) violations of the conservativity principle.

only occur over the existing set of correspondences $\mathcal{A}$ and the candidate correspondence c . Note that each agent has only full access to its own ontology (i.e., $\mathcal{O}^x$) and the ontology of the other agent is seen as private (i.e., $\mathcal{O}^{\hat{x}}$). A repair, given $\mathcal{A}$ and a candidate correspondence c , is a set of correspondences whose removal from the alignment would eliminate all violations. We define *Alignment Repair* as follows:

Definition 6: Let $\mathcal{A}'$ be the new set of correspondences $\mathcal{A} \cup \{c\}$, where c is a candidate correspondence and $\mathcal{A}$ is the current alignment w.r.t. $\mathcal{O}^x$, for which there is a consistency or conservativity violation. An alignment $\mathcal{R} \subseteq \mathcal{A}'$ is a repair for $\mathcal{A}'$ w.r.t. $\mathcal{O}^x$ iff there are no such violations in $\mathcal{O}^x \cup \mathcal{A}' \setminus \mathcal{R}$.

A trivial repair is $\mathcal{R} = \{c\}$, as the removal of the candidate correspondence c that introduces a violation would obviously eliminate that violation. However, the objective is to remove as little (useful) information as possible (i.e., a repair may result in the weakening of existing equivalence correspondences⁴). Furthermore, in case of multiple options, the correspondence weight will be used as a differentiating factor (i.e., a correspondence with a lower weight will be weakened over a correspondence with higher weight). When a correspondence is weakened, it “inherits” its original weight; so that this can be considered for future repairs.

Agents *rationaly* determine whether correspondences causing violations should be repaired or rejected. The strategy we employ favours *stability* over *maximality*, i.e., higher weighted correspondences are preferred over lower weighted ones (even if the cumulative aggregate weight of the repairs is greater than that of the higher weighted correspondence currently proposed). Hence, repair is considered *rational* if the proposed correspondence has a higher weight

⁴As an equivalence correspondence ($a \equiv b$) $\models (a \sqsubseteq b) \sqcap (b \sqsubseteq a)$, it can be weakened by eliminating one of the two subsumptions.

Table 2: The private and joint *weights* for the correspondences in the worked example, and the two ontology fragments assumed by the two agents *Alice* and *Bob*. Only *Alice* is aware of the correspondence $\langle d, y, \equiv \rangle$.

<i>Alice</i> $b \sqsubseteq a$ $c \sqsubseteq a$ $b \sqcap c \sqsubseteq \perp$ $c \equiv d$ Private Ontology		<i>Bob</i> $x \equiv y$ $z \sqsubseteq y$ w Private Ontology	
Public Signature a b c		Public Signature w x y z	
c	κ_c^{Alice}	κ_c^{Bob}	κ_c^{joint}
$\langle a, w, \equiv \rangle$	0.35	0.25	0.3
$\langle a, x, \equiv \rangle$	0.9	0.8	0.85
$\langle b, x, \equiv \rangle$	0.6	0.6	0.6
$\langle b, y, \equiv \rangle$	0.4	0.7	0.55
$\langle b, z, \equiv \rangle$	0.55	0.6	0.575
$\langle c, y, \equiv \rangle$	0.25	0.65	0.45
$\langle d, y, \equiv \rangle$	0.7	—	—

than all of those proposed for deletion by the repair plan. A repair is *not rational* if the inclusion of the proposed correspondence results in the deletion of a higher weighted correspondence.

Correspondence Repair Dialogue Example

In this section we illustrate by means of an example how agents engage in the dialogue and the strategic decisions they make when assessing whether or not a correspondence should be included in the final alignment. The example shows how the agents can repair correspondences they need to propose (or that they receive) and that may cause conservativity and/or consistency violations.

Two agents, *Alice* and *Bob*, each possess a private ontological fragment, that provides the conceptualisation for the entities that they use to communicate (Table 2). Each agent has acquired a subset of correspondences, with an associated weight κ_c in the range $[0, 1]$, which is initially private to each agent. These are summarised (with the resulting joint(c) for each c) in Table 2. Finally, both agents assume that the *admissibility threshold* $\epsilon = 0.45$ to filter out correspondences with a low joint(c).

The example dialogue between *Alice* and *Bob* is presented in Table 3. The two agents initiate the dialogue by both uttering the *join* move (Moves 1-2, omitted from Table 3), and the turn order is non-deterministic; in this example, *Alice* makes the first *join* move. Each exchange is shown with its move identifier, and the state (taken from Figure 1) from which the move is taken.

Move 3: *Alice* selects one of her undisclosed correspondences with the highest κ_c ; in this case, $\langle a, x, \equiv \rangle$. Initially, *Alice* assumes *Bob*'s upper bound $\kappa_u^{Bob} = 1$, and estimates the joint weight for c , $\kappa_{\langle a, x, \equiv \rangle}^{est} = \frac{1}{2}(0.9 + 1) = 0.95$. As this is equal to or above threshold ($\epsilon = 0.45$), she asserts c . Given that $\mathcal{A} = \emptyset$, the inclusion of c introduces no violation and thus, no repair is necessary.

Move 4: *Bob* confirms that the joint weight $\kappa_{\langle a, x, \equiv \rangle}^{joint} =$

Table 3: The messages exchanged between *Alice* and *Bob* in the example dialogue. We omit the two initial *join* moves for sake of brevity and we start from move 3.

Move	State	Locution
3	S3	$\langle Alice, assert, \langle a, x, \equiv \rangle, 0.9, \emptyset \rangle$
4	S4	$\langle Bob, acceptC, \langle a, x, \equiv \rangle, 0.85, nil \rangle$
5	S6	$\langle Bob, assert, \langle b, y, \equiv \rangle, 0.7, \emptyset \rangle$
6	S7	$\langle Alice, acceptC, \langle b, y, \equiv \rangle, 0.55, nil \rangle$
7	S3	$\langle Alice, assert, \langle b, x, \equiv \rangle, 0.6, \{ \langle b, x, \sqsupset \rangle \} \rangle$
8	S4	$\langle Bob, acceptC, \langle b, x, \sqsupset \rangle, 0.6, nil \rangle$
9	S6	$\langle Bob, assert, \langle c, y, \equiv \rangle, 0.65, \emptyset \rangle$
10	S7	$\langle Alice, rejectC, \langle c, y, \equiv \rangle, nil, nil \rangle$
11	S3	$\langle Alice, assert, \langle b, z, \equiv \rangle, 0.55, \emptyset \rangle$
12	S4	$\langle Bob, repair, \langle b, z, \equiv \rangle, 0.575, \{ \langle b, y, \sqsupset \rangle \} \rangle$
13	S5	$\langle Alice, acceptR, \langle b, z, \equiv \rangle, nil, nil \rangle$
14	S6	$\langle Bob, close, nil, nil, nil \rangle$
15	S1	$\langle Alice, close, nil, nil, nil \rangle$

$\frac{1}{2}(0.9 + 0.8) = 0.85$ is above threshold, and checks to see if any repair is needed before accepting the correspondence. As none is needed, he simply accepts c , and notifies *Alice* of the joint weight $\kappa_{\langle a, x, \equiv \rangle}^{joint}$. He adds c to $\mathcal{A}$ and updates *Alice*'s upper bound $\kappa_u^{Alice} = 0.9$; as this was *Alice*'s highest weighted correspondence, she will have no other undisclosed c where $\kappa_c^{Alice} > \kappa_u^{Alice}$. On receipt of this move, *Alice* adds c to $\mathcal{A}$.

Moves 5-6: *Bob* selects his highest privately-weighted undisclosed correspondence $c = \langle b, y, \equiv \rangle$. He estimates the joint weight $\kappa_{\langle b, y, \equiv \rangle}^{est} = \frac{1}{2}(0.7 + \kappa_u^{Alice}) = 0.55$, which is above threshold, and finds that no repairs are necessary if c is added to $\mathcal{A}$. He therefore asserts c . *Alice* confirms that $\kappa_{\langle b, y, \equiv \rangle}^{joint} = \frac{1}{2}(0.7 + 0.4) = 0.55$ is above threshold, and from her perspective, no repairs are necessary. She accepts the correspondence, adds c to $\mathcal{A}$ and updates her upper bound $\kappa_u^{Bob} = 0.7$.

At this point, both agents have the alignment $\mathcal{A} = \{ \langle a, x, \equiv \rangle, \langle b, y, \equiv \rangle \}$. If we consider $\mathcal{O}^{Alice} \cup \mathcal{A} \cup \mathcal{O}^{Bob}$ then we have potential violations: given *Bob*'s axiom $x \equiv y$, then $\mathcal{A} \cup \{ x \equiv y \} \models (b \sqsubseteq a) \wedge (a \sqsubseteq b)$. However, as *Bob* has only the axiom $b \sqsubseteq a$, adding $\{ x \equiv y \}$ to $\mathcal{A}$ would introduce a new axiom, thus violating conservativity. Additionally, $\mathcal{A} \cup \{ x \equiv y \} \models (c \sqsubseteq b)$ also holds, thus there is a consistency violation due to the axiom $(b \sqcap c \sqsubseteq \perp)$ in $\mathcal{O}^{Alice}$. It is important to note that these violations only occur if both ontologies are known, which is not the case in the current example. If we consider each ontology individually with $\mathcal{A}$, then no violations occur. In the next move, we consider the case where a violation is introduced given a single ontology.

Moves 7-8: *Alice* had previously acquired the correspondence $\langle d, y, \equiv \rangle$ from an earlier encounter. However, as it maps an entity (d) that she does not want to reveal to *Bob* (and thus is not a member of her public signature), she identifies the next viable correspondence, $\langle b, x, \equiv \rangle$, which is above threshold; i.e., $\kappa_{\langle b, x, \equiv \rangle}^{est} = \frac{1}{2}(0.6 + \kappa_u^{Bob}) = 0.6$. However, the inclusion of this correspondence would introduce a conservativity violation for *Alice*, as $\mathcal{O}^{Alice} \cup \mathcal{A} \cup \mathcal{O}^{Bob} \models (a \sqsubseteq b) \wedge (b \sqsubseteq a)$, but $\mathcal{O}^{Alice}$ contains only

Table 4: The status of the dialogue after Move 10.

Alice's Ontology	Alignment	Bob's Ontology
$\begin{array}{c} a \\ \sqsubseteq \quad \sqsubseteq \\ b \quad c \equiv d \end{array}$	$\begin{array}{c} w \\ \sqsubseteq \quad \sqsubseteq \\ a \equiv x \\ b \sqsubseteq y \\ c \equiv z \end{array}$	$\begin{array}{c} x \equiv y \\ \sqsubseteq \\ w \quad z \\ x \equiv y \\ z \sqsubseteq y \end{array}$
$\begin{array}{ll} b \sqsubseteq a & b \sqcap c \sqsubseteq \perp \\ c \sqsubseteq a & c \equiv d \end{array}$		
$\kappa_u^{Alice} = 0.6$		
$\kappa_u^{Bob} = 0.65$		
$\mathcal{A} = \{\langle a, x, \equiv \rangle, \langle b, x, \sqsubseteq \rangle, \langle b, y, \equiv \rangle\}$		
$c \in CS = \{\langle a, x, \equiv \rangle, \langle b, y, \equiv \rangle, \langle b, x, \equiv \rangle, \langle c, y, \equiv \rangle\}$		
$c \notin CS = \{\langle b, z, \equiv \rangle, \langle a, w, \equiv \rangle, \langle d, y, \equiv \rangle\}$		

$b \sqsubseteq a$. We assume that the ontologies cannot be changed, and if a violation is detected, the only rational move for an agent is propose a repair for the violation detected by either removing one or more correspondences from the alignment computed, $\mathcal{A}$ or by accepting a weakened version of the proposed correspondence, based on the correspondences utilities. In this case, *Alice* can either repair $\langle a, x, \equiv, 0.85 \rangle$ or $\langle b, x, \equiv, 0.5 \rangle$. She retains $\langle a, x, \equiv, 0.85 \rangle$ as it has the higher joint weight, and removes the axiom $\langle b, x, \sqsubseteq, 0.5 \rangle$ from $\mathcal{A}$. As the repair does not cause any violation to *Bob*, he accepts the repair, then updates the upper bound $\kappa_u^{Alice} = 0.6$.

Moves 9-10: *Bob* identifies the next viable correspondence: $\langle c, y, \equiv \rangle$, which is above threshold given *Alice*'s upper bound $\kappa_u^{Alice} = 0.6$; i.e., $\kappa_{\langle c, y, \equiv \rangle}^{est} = \frac{1}{2}(0.65 + \kappa_u^{Alice}) = 0.625$. *Alice*, however, detects a *consistency* violation caused by the inclusion of this correspondence in $\mathcal{A}$, since $\mathcal{A} \models (b \equiv c)$, but $\mathcal{O}^{Alice}$ contains the axiom $(b \sqcap c \sqsubseteq \perp)$, thus making $b \equiv c$ unsatisfiable. *Alice* can propose two possible repairs: 1) reject $c \equiv y$; or 2) remove $b \equiv y$ and include $c \equiv y$ to the commitment store. However, this second choice would not guarantee that further repairs could be avoided. As $\kappa_{\langle b, y, \equiv \rangle}^{joint} > \kappa_{\langle c, y, \equiv \rangle}^{joint}$ (see Table 2), *Alice* rejects $\langle c, y, \equiv \rangle$. The status of the dialogue at this point is illustrated in Table 4.

Moves 11-13: *Alice* asserts $\langle b, z, \equiv \rangle$ as $\kappa_{\langle b, z, \equiv \rangle}^{joint} = 0.575$.

When *Bob* assesses this correspondence, it detects a conservativity violation, since his ontology contains the axiom $z \sqsubseteq y$, and the inclusion of both $b \equiv y$ and $b \equiv z$ would also infer the axiom $y \sqsubseteq z$ (similarly for $x \sqsubseteq z$). As $\kappa_{\langle b, y, \equiv \rangle}^{joint} < \kappa_{\langle b, z, \equiv \rangle}^{joint}$, *Bob* suggests a repair that weakens $b \equiv y$ by removing $b \sqsubseteq y$ and thus leaving the correspondence $b \sqsubseteq y$. *Alice* confirms that the non-weakened version of the asserted correspondence $\kappa_{\langle b, z, \equiv \rangle}^{joint} = 0.575 \geq \epsilon$, and that no further violations are detected as a consequence of the repair, hence she accepts the assertion. *Bob* updates his upper estimate of *Alice*'s weight, $\kappa_u^{Alice} = 0.55$.

Moves 14-15: *Bob* can now estimates the joint utility of his remaining correspondence, $\langle a, w, \equiv \rangle$, but $\kappa_{\langle a, w, \equiv \rangle}^{est} =$

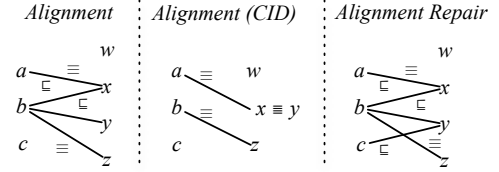


Figure 2: The final alignment $\mathcal{A}$, with the two ontologies. The second alignment (Alignment CID) represents the result if the example had been evaluated using the original *Correspondence Inclusion Dialogue*. The third alignment is the one generated by using only the repair mechanism.

$\frac{1}{2}(0.25 + \kappa_u^{Alice}) = 0.4 < \epsilon$. *Bob* therefore realises that he has no further correspondence to propose and issues a *close* move. *Alice* reduces her upper upper estimate of *Bob*'s weight, $\kappa_u^{Bob} = \epsilon$, due to the *close* move. Although she has a remaining correspondence to propose $\langle a, w, \equiv \rangle$, she estimates its joint weight $\kappa_{\langle a, w, \equiv \rangle}^{est} = \frac{1}{2}(0.35 + \kappa_u^{Bob}) = 0.4 < \epsilon$. As this is below threshold, and she has no other viable correspondences to assert, she utters a *close* move, and the dialogue terminates.

The resulting alignment $\mathcal{A} = \{\langle a, x, \equiv \rangle, \langle b, x, \sqsubseteq \rangle, \langle b, y, \sqsubseteq \rangle, \langle b, z, \equiv \rangle\}$ is illustrated in Figure 2, against the alignments that would be generated by the original *Correspondence Inclusion Dialogue* (Payne and Tamma 2014b) and by the original repair mechanism (Solimando, Jiménez-Ruiz, and Guerrini 2014b). When compared to the dialogue with no repairs, the solution presented here results in two further correspondences, $\langle b, x, \sqsubseteq \rangle$ and $\langle b, y, \sqsubseteq \rangle$, both of which are consistent with the fact that: 1) from *Alice*'s perspective she has the correspondence $\langle a, x, \equiv \rangle$ and that $b \sqsubseteq a$; or 2) from *Bob*'s perspective he has the correspondence $\langle b, z, \equiv \rangle$ and that $z \sqsubseteq y$. The solution found by the repair system contains an additional correspondence, $\langle c, y, \sqsubseteq \rangle$, that our decentralised mechanism rejected as a possible repair because it was less preferable than an alternative repair plan (Moves 9-10). This is an artefact of the fact that the *Correspondence Inclusion Dialogue* is intrinsically suboptimal because both agents must agree on each correspondence (i.e., the joint needs to be above the admissibility threshold). Furthermore, given the assumption of partial knowledge, it is not always the case that agents find a repair and correspondences might be rejected because they can potentially cause (local) violation that the other agent cannot detect, and for which it would not suggest a repair.

Dialogue Properties

The dialogue mechanism presented in the previous section is not meant to be complete, and it produces a sub-optimal solution to the alignment problem. By relaxing the assumptions that both ontologies are shared, and by computing the consequences of the union of the individual ontologies and the alignment, the solutions found are incomplete. However, the aim of the approach presented here is to assess whether some repairs can be unilaterally computed and if the detec-

tion of violations can be included in the agents strategies; whereby they do not select correspondences that cause violations, or propose repairs if these need to be included. While we do not claim that our approach is complete, it is possible to prove that the dialogues terminates.

Proposition 1 *The dialogue with the set of moves $\mathcal{M}$ in Table 1 will always terminate.*

Proof Both the agents engaging in the dialogue, the *sender* x and the *receiver* $\hat{x}$ have finite correspondence stores Δ^x and $\Delta^{\hat{x}}$, as the set of possible correspondences $\mathcal{C}$ is finite. In the dialogue the agents can propose a correspondence only once; when a correspondence is disclosed, it is added to the commitment stores (that are always kept synchronised). The dialogue consists of a sequence of *interactions*, i.e., moves delimited by *assert* – *close* that determines the inclusion of a correspondence c . When c is asserted, it is either accepted or rejected, depending on its assessment with respect to the *admissibility threshold* and the *repair* strategy. However, once a correspondence has been disclosed it can no longer be retracted. If the dialogue does not end before every possible correspondence is considered for inclusion, then it will end once the (finite) set of *assert* moves regarding the possible correspondences have all been made once. This is contained in $\mathcal{C}$ as we only assert correspondences that are viable and above threshold.

Regarding the repair strategy, this is a self contained module that is invoked by the dialogue mechanisms. More information about the repair techniques can be found in (Solimando, Jiménez-Ruiz, and Guerrini 2014b). Its termination follows trivially from the techniques used: detecting subsumption violations translates to visiting a finite labelled graph, following (Dowling and Gallier 1984), while detecting the equivalence violations requires the agents to compute the strongly connected components (SCCs) of a graph representation of the ontologies, using Tarjan’s algorithm (Tarjan 1972). Concerning the repair technique for subsumption violations, it requires the agents to explore the different subsets of the mappings involved in at least one violation. The repair of equivalence violations is computed by selecting a subset of the arcs in each of the SCCs presenting at least one violation. This task has a number of steps bounded by the powerset of the arcs of each SCC in the graph. ■

As already discussed in the Preliminaries section, the repair mechanism proposed by Jiménez-Ruiz et al. (2011) is *sound* with the respect to the violations that would be detected by a centralised approach or by full reasoning. The soundness property is a direct consequence of the soundness of the algorithms used for the detection of cycles in the graph representation of the ontologies and the alignments. The dialogue presented in this paper does not affect soundness as it only determines whether a repair to a detected violation is to be proposed or rejected, but does not alter a repair plan, as repairs can be only accepted or rejected.

Empirical Evaluation

The purpose of the empirical evaluation is to assess the impact of incompleteness on the performance of the dialogue.

As discussed in the previous section, the dialogue is guaranteed to terminate and it is sound, however it is incomplete by its very nature as it focusses on sub-optimal solutions. The aim of this evaluation is to verify that even if we do not generate the whole set of solutions, the dialogue is still fit for purpose, i.e., that it can find solutions comparable with the baseline provided by state of the art alignment systems. The empirical analysis of the effectiveness of the proposed dialogue is assessed in terms of number of mappings exposed, number of mappings agreed, precision and recall (with respect to a reference alignment). We report these metrics with respect to different values of the admissibility threshold ϵ .

The following two hypotheses have been tested using *Ontology Evaluation Alignment Initiative* (OAEI) data sets:⁵

1. Selecting and combining correspondences taken from several different alignment methods can yield comparable performance to existing alignment methods, when measured using the *precision*, *recall* and *f-measure* metrics;
2. Eliminating low utility correspondences improves dialogue performance with respect to the resulting alignment, and the number of correspondences disclosed;

The OAEI 2012 Conference Track⁶ comprises various ontologies describing the same conceptual domain (conference organisation) and alignments between pairs of ontologies, generated by 18 different ontology alignment approaches. Seven ontologies were selected as these were accompanied by *reference alignments* (defined by a domain expert), resulting in 18 different alignments for each pair of ontologies, and $\frac{7!}{(7-2)!2!} = 21$ ontology combination pairs.

The empirical evaluations were conducted over each of the 21 ontology pairs (i.e., for each experiment, the two agents were allocated a pair of ontologies (one each) and a set of alignments previously generated by the systems participating to the 2012 OAEI competition. From these different alignments, each agent was able to generate a probability distribution over the alignments, such that $p(c)$ reflected the frequency of the occurrence of the correspondence c in the agent’s available alignments. This is based on the assumption that a correspondence found in many alignments has a higher likelihood to be correct. In this evaluation, we relax the asymmetric knowledge assumption by assigning both agents the same 18 alignments to facilitate verification of the resulting dialogue. Experiments were also repeated for different admissibility thresholds to evaluate how it could affect the number of messages exchanged, and consequently the correspondences disclosed.

The resulting alignments were evaluated using the *precision*, *recall* and *f-measure* metrics, where: **precision** (p) is the proportion of correspondences found by the dialogue that are correct (i.e., in the *reference alignment*); **recall** (r) is the proportion of correct correspondences w.r.t. the number of correspondences in the *reference alignment*; and the **f-measure** (f) represents the harmonic mean of p and r .

⁵<http://oaei.ontologymatching.org>

⁶Alignments have also been tested from more recent OAEI competitions, but with no discernible difference in the results.

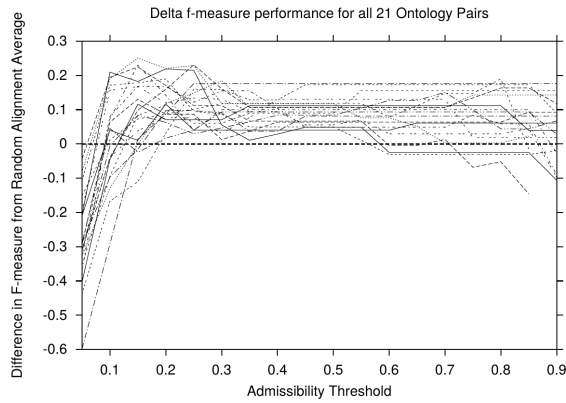


Figure 3: Delta f-measures (δf_D) for all 21 ontology pairs.

A baseline was generated by assuming that a naïve approach for finding an alignment would consist of an agent randomly picking and using one of the pre-defined alignments. Thus, we compute the average performance of the 18 alignment methods for each ontology pair, and use this to evaluate the comparative performance of the alignments generated by the dialogue.

Figure 3 demonstrates how, in most cases, the f-measure performance of the dialogue is significantly higher than that achieved from selecting an alignment at random, when (in most cases) $\epsilon > 0.16$, *i.e.*, correspondences are only considered if they appeared in at least three of the original OAEI alignments. The graph in Figure 3 plots the difference in f-measure (denoted δf_D) between that achieved using the average alignments ($\bar{f}_A$), and that achieved by the dialogue for all 21 ontology pairs (*i.e.*, values above zero indicate a better f-measure, whereas those below are worse). As ϵ increases, the rise in precision generated for each ontology pair stabilises, whereas the recall starts to fall due to viable correspondences being rejected. This results in a corresponding decline in the f-measure, which starts to affect some ontology pairs at $\epsilon > 0.6$, and is increasingly noticeable as $\epsilon > 0.85$.

Related Work

A number of different approaches have addressed the reconciliation of heterogeneous ontologies by using some form of rational reasoning. Argumentation has been used as a rational means for agents to select ontology correspondences based on the notion of partial-order preferences over their different properties (*e.g.*, structural vs terminological) (Laera et al. 2007). A variant was proposed by Trojahn dos Santos, Quaresma, and Vieira (2008) which represented ontology mappings as disjunctive queries in Description Logics. Typically, these approaches have used a coarse-grained decision metric based on the *type* of correspondence, rather its *acceptability* to each agent (given other mutually accepted correspondences), and do not consider the notion of private, or asymmetric knowledge, but assume the correspondences to be publicly accessible.

The conservativity and consistency principles have been exploited by several ontology alignment systems, *e.g.*,

ASMOV (Jean-Mary, Shironoshita, and Kabuka 2009), *Lily* (Wang and Xu 2008) and *YAM++* (Ngo and Bellahsene 2012) implemented different heuristics in order to detect conservativity violations and improve precision with respect to a reference alignment. (Beisswanger and Hahn 2012), in contrast, proposes a set of sanity checks and best practices to use when computing the mappings. In this paper we have reused the detection and repair mechanism proposed by (Solimando, Jiménez-Ruiz, and Guerrini 2014a). However, alternative repair systems, such as *ALCOMO* (Meilicke 2011) could also have been used. Our approach shows how agents use strategic decision making to choose whether or not to repair the alignment being computed. Crucially, the repair is an integral part of the alignment generation, and considers the consequences of the partially generated alignment when it is integrated with the agent’s ontology. This allows the repair mechanism to propose the removal of all those correspondences that cause a violation. Earlier work (Dos Santos and Euzenat 2010) focussed on detecting inconsistency caused by two incompatible correspondences and could not handle the consequences of several axioms.

Other approaches like (Lambrix and Liu 2013) consider conservativity violations as false positives, and hence their correction strategy aims at extending the is-a relationships of the input ontologies. (Pesquita et al. 2013) also question the generation of (alignment) repairs and suggest to change the ontologies. Our repair strategy, however, follows a “better safe than sorry” approach, suitable for the scenario presented in this paper where the ontologies are not modifiable.

Conclusions

This paper presents a novel inquiry dialogue that significantly extends the *Correspondence Inclusion Dialogue* described in (Payne and Tamma 2014a; 2014b). The dialogue facilitates negotiation over asymmetric and incomplete knowledge of ontological correspondences. It enables two agents to selectively disclose private correspondences given their perceived utility. Correspondences are only permitted when they do not introduce consistency or conservativity violations for each agent’s ontology in isolation. A running example is presented, that illustrates how the conservativity violation repairs are shared and applied. An implementation of the dialogue and repair mechanism have been used to evaluate the efficacy of the approach for the negotiation of correspondences taken from the OAEI test data.

As immediate future work we plan to provide a more fine-grained evaluation where agents may opt to use different repair strategies (*e.g.*, only consistency repair or conservativity violations repair). We also aim at extending the dialogue among agents by allowing to disclose some additional knowledge about their (private) ontology. The partial disclosure of the agents ontology will enable the integration of the *locality* principle (Jiménez-Ruiz et al. 2011) within the dialogue, and enhance the repair decisions for the already integrated consistency and conservativity principles. Furthermore we also plan to extend the dialogue to consider scenarios involving more than two agents. In such scenarios new (repair) challenges arise (Euzenat 2015) and our repair techniques will most likely need to be adapted.

Acknowledgments

This work was funded by the EU project Optique (FP7-ICT-318338) and the EPSRC projects DBOnto, ED3 and MaSI³. The authors would like to thank the anonymous reviewers and Bijan Parsia for their insightful comments on previous versions of this paper.

References

- Beisswanger, E., and Hahn, U. 2012. Towards valid and reusable reference alignments - ten basic quality checks for ontology alignments and their application to three different reference data sets. *J. Biomed. Sem.* 3(S-1):S4.
- Borgida, A., and Serafini, L. 2003. Distributed Description Logics: Assimilating Information from Peer Sources. *J. Data Sem.* 1:153–184.
- Cheatham, M., et al. 2015. Results of the Ontology Alignment Evaluation Initiative 2015. In *OM Workshop*, 60–115.
- Cuenca Grau, B., and Motik, B. 2012. Reasoning over ontologies with hidden content: The import-by-query approach. *J. Artif. Intell. Res. (JAIR)* 45.
- Doran, P.; Tamma, V.; Payne, T. R.; and Palmisano, I. 2009. Dynamic selection of ontological alignments: a space reduction mechanism. In *Int'l Joint Conf. Artif. Intell. (IJCAI)*.
- Dos Santos, C. T., and Euzenat, J. 2010. Consistency-driven argumentation for alignment agreement. In *OM Workshop*, 37–48.
- Dowling, W. F., and Gallier, J. H. 1984. Linear-Time Algorithms for Testing the Satisfiability of Propositional Horn Formulae. *J. Log. Program.* 1(3):267–284.
- Euzenat, J., and Shvaiko, P. 2013. *Ontology Matching, Second Edition*. Springer.
- Euzenat, J. 2015. Revision in networks of ontologies. *Artif. Intell.* 228:195–216.
- Gale, D., and Shapley, L. S. 1962. College admissions and the stability of marriage. *Amer. Math. Monthly* 69(1):9–15.
- Grice, H. P. 1975. Logic and conversation. In *Syntax and Semantics: Vol. 3: Speech Acts*. Academic Press. 41–58.
- Jean-Mary, Y. R.; Shironoshita, E. P.; and Kabuka, M. R. 2009. Ontology matching with semantic verification. *J. Web Sem.* 7(3).
- Jiménez-Ruiz, E., and Cuenca Grau, B. 2011. LogMap: Logic-based and scalable ontology matching. In *Int'l Sem. Web Conf. (ISWC)*, 273–288.
- Jiménez-Ruiz, E.; Cuenca Grau, B.; Horrocks, I.; and Berlanga, R. 2009. Ontology integration using mappings: Towards getting the right logical consequences. In *Eur. Sem. Web Conf. (ESWC)*, 173–187.
- Jiménez-Ruiz, E.; Cuenca Grau, B.; Horrocks, I.; and Berlanga, R. 2011. Logic-based assessment of the compatibility of UMLS ontology sources. *J. Biomed. Sem.* 2:S2.
- Kalyanpur, A.; Parsia, B.; Horridge, M.; and Sirin, E. 2007. Finding all justifications of OWL DL entailments. In *Int'l Sem. Web Conf. (ISWC)*, 267–280.
- Konev, B.; Walther, D.; and Wolter, F. 2008. The Logical Difference Problem for Description Logic Terminologies. In *Int'l Joint Conf. on Autom. Reason. (IJCAR)*, 259–274.
- Kuhn, H. W. 1955. The Hungarian method for the assignment problem. *Naval research logistics quarterly* 2(1-2).
- Laera, L.; Blacoe, I.; Tamma, V.; Payne, T.; Euzenat, J.; and Bench-Capon, T. 2007. Argumentation over ontology correspondences in MAS. In *Int'l Conf. on Auton. Agent and Multi-Agent Syst. (AAMAS)*, 1285–1292.
- Lambrix, P., and Liu, Q. 2013. Debugging the Missing Is-a Structure Within Taxonomies Networked by Partial Reference Alignments. *Data Knowl. Eng. (DKE)* 86:179–205.
- Meilicke, C. 2011. *Alignment Incoherence in Ontology Matching*. Ph.D. Dissertation, Dissertation, Universität Mannheim, Mannheim.
- Ngo, D., and Bellahsene, Z. 2012. YAM++ : A Multi-strategy Based Approach for Ontology Matching Task. In *Int'l Knowl. Eng. and Knowl. Management Conf. (EKAW)*.
- Payne, T. R., and Tamma, V. 2014a. A dialectical approach to selectively reusing ontological correspondences. In *Int'l Knowl. Eng. and Knowl. Management Conf. (EKAW)*.
- Payne, T. R., and Tamma, V. 2014b. Negotiating over ontological correspondences with asymmetric and incomplete knowledge. In *Int'l Conf. on Auton. Agent and Multi-Agent Syst. (AAMAS)*, 517–524.
- Pesquita, C.; Faria, D.; Santos, E.; and Couto, F. M. 2013. To repair or not to repair: reconciling correctness and coherence in ontology reference alignments. In *OM Workshop*, 13–24.
- Schlobach, S. 2005. Debugging and Semantic Clarification by Pinpointing. In *Eur. Sem. Web Conf. (ESWC)*, 226–240.
- Shvaiko, P., and Euzenat, J. 2013. Ontology matching: State of the art and future challenges. *IEEE Trans. Knowl. Data Eng.* 25(1):158–176.
- Solimando, A.; Jiménez-Ruiz, E.; and Guerrini, G. 2014a. A Multi-strategy Approach for Detecting and Correcting Conservativity Principle Violations in Ontology Alignments. In *OWLED Workshop*, 13–24.
- Solimando, A.; Jiménez-Ruiz, E.; and Guerrini, G. 2014b. Detecting and correcting conservativity principle violations in ontology-to-ontology mappings. In *Int'l Sem. Web Conf. (ISWC)*, 1–16.
- Solimando, A.; Jiménez-Ruiz, E.; and Guerrini, G. 2015. On the feasibility of using OWL 2 reasoners in ontology alignment repair problems. In *ORE Workshop*, 60–67.
- Tarjan, R. E. 1972. Depth-first search and linear graph algorithms. *SIAM J. Comput.* 1(2):146–160.
- Trojahn dos Santos, C.; Quaresma, P.; and Vieira, R. 2008. Conjunctive queries for ontology based agent communication in MAS. In *Int'l Conf. on Auton. Agent and Multi-Agent Syst. (AAMAS)*, 829–836.
- Walton, D., and Krabbe, E. 1995. *Commitment in Dialogue: Basic Concepts of Interpersonal Reasoning*. SUNY series in Logic and Language. State University of New York Press.
- Wang, P., and Xu, B. 2008. Debugging ontology mappings: A static approach. *Comput. Inform.* 27(1):21–36.