

Dialectical Characterization of Consistent Query Explanation with Existential Rules

Abdallah Arioua^{1 2} and Madalina Croitoru¹

¹ GraphIK, Univ. of Montpellier, France

² UMR IATE - INRA , Montpellier, France

Abstract

This paper provides a dialectical characterization of $\mathcal{IAR}$ and $\mathcal{Brave}$ semantics via argumentation dialogue. We propose a minimal argumentation dialogue system and investigate the relation between its outcome and these semantics. We show that $\mathcal{Brave}$ semantics corresponds to a dialogue where the opponent of the $\mathcal{Brave}$ -entailed query wins in any dialogue. We show also that $\mathcal{IAR}$ semantics corresponds to a dialogue where the proponent of the $\mathcal{IAR}$ -entailed query wins any dialogue. We further investigate how the profile (i.e. behavior) of the participants impacts the outcome of the dialogue and the entailment of queries.

Introduction

Many recent works in knowledge representation have focused on addressing the problem of consistent query answering within expressive logical languages such as the existential rule setting (Baget et al. 2011b; Bienvenu 2012; Lembo et al. 2010). The problem consists of investigating the logical and computational properties of various inconsistency-tolerant semantics. Inconsistency-tolerant semantics consider as input a set of factual knowledge bases, each of which is consistent with an ontology, while their union (the merger of the knowledge bases) is inconsistent. The semantics consider the repairs, i.e. the maximal consistent subsets of the union of knowledge bases and provide different recipes for using these repairs when answering queries. For instance the $\mathcal{Brave}$ semantics considers enough the fact that *at least* one repair entails the query. $\mathcal{IAR}$ semantics is far more restrictive: the query must be entailed from the intersection of all repairs. Other semantics have been considered between these two extremes but here we will focus only on $\mathcal{Brave}$ and $\mathcal{IAR}$.

One of the major shortcomings (up to recent works (Bienvenu, Bourgaux, and Goasdoué 2016; Arioua, Tamani, and Croitoru 2015)) is the lack of explanation facilities for inconsistency-tolerant semantics.

(Bienvenu, Bourgaux, and Goasdoué 2016) started considering inconsistency-tolerant semantics explanation based on the notion of *causes*. They mainly focus on computational aspects of finding explanations and do not cover the representational aspect with the user.

Copyright © 2016, Association for the Advancement of Artificial Intelligence (www.aaai.org). All rights reserved.

(Arioua, Tamani, and Croitoru 2015; Arioua et al. 2014) consider inconsistency-tolerant semantics explanation with another framework, i.e. logical argumentation. However the authors only consider the $\mathcal{ICR}$ semantics (Bienvenu 2012), in addition no necessary and sufficient conditions are given to characterize an argumentation dialogue with respect to $\mathcal{ICR}$ semantics.

In the light of the state of the art, in this paper we provide a dialectical characterization of the $\mathcal{Brave}$ and $\mathcal{IAR}$ semantics. More precisely we would like to be able to give necessary and sufficient conditions a dialogue should satisfy with respect to the $\mathcal{Brave}$ and $\mathcal{IAR}$ semantics.

In order to achieve this we first propose an argumentation dialogue system that considers a turn taking game between a proponent and an opponent. Similar to (Parsons, Wooldridge, and Amgoud 2003) we define the concept of participant's profile and depending on these profiles we will be able to give necessary and sufficient conditions for the $\mathcal{Brave}$ and $\mathcal{IAR}$ semantics.

The paper is structured as follows. Section 1 introduces the logical language used in this paper alongside with the inconsistency-tolerant semantics. Section 2 presents the logic-based argumentation framework. Next, Section 3 & 4 present the contribution of the paper where we introduce the argumentation dialogue system A_0 and the dialectical characterization of the $\mathcal{IAR}$ and $\mathcal{Brave}$ semantics. Finally, Section 5 concludes the paper.

Logical Language

We consider the existential rules fragment of first-order logic which is composed of formulas built with the only logical connectives ($\wedge$, $\rightarrow$) (conjunction and implication), the quantifiers ($\exists$, $\forall$) and the special constant $\perp$ (falsum). An *atom* is of the form $p(t_1, \dots, t_k)$ where p is a predicate of arity k and the t_i are terms, i.e. variables or constants. A finite set of atoms F is called an *atomset*, we denote by $terms(F)$ (resp. $vars(F)$) the set of terms (resp. variables) that occur in F . Given atomsets A_1 and A_2 , a *homomorphism* π from A_1 to A_2 is a substitution of $vars(A_1)$ by $terms(A_2)$ such that $\pi(A_1) \subseteq A_2$.

An *existential rule* (or a rule) is of the form $R = \forall \vec{x} \forall \vec{y} (B \rightarrow \exists \vec{z} H)$, where B and H are conjunctions of atoms, with $vars(B) = \vec{x} \cup \vec{y}$, and $vars(H) = \vec{x} \cup \vec{z}$. B and H are respectively called the *body* and the *head* of R .

A rule with an empty body (resp. head set to $\perp$) is called a *fact* (resp. *negative constraint*). A *Boolean conjunctive query* (BCQ) has the form of a fact. From now on we use the general term *query* to mean BCQ.

A *knowledge base* $\mathcal{K} = (\mathcal{F}, \mathcal{R}, \mathcal{N})$ is composed of a finite sets of facts $\mathcal{F}$, rules $\mathcal{R}$ and negative constraints $\mathcal{N}$. In the following, we see conjunctions of atoms as atomsets. A rule $R : B \rightarrow H$ is *applicable* to an atomset F if there is a homomorphism π from B to F . The *application of R to F* w.r.t. π produces an atomset $\alpha(F, R, \pi) = F \cup \pi(\text{safe}(H))$, where $\text{safe}(H)$ is obtained from H by replacing existential variables with fresh variables. Using the *chase* we deduce new facts by applying rules on the initial set of facts $\mathcal{F}$. We restrict our work to the *finite expansion set* of rules $\mathcal{R}$ on which the chase always **halts** for any atomset F (Baget et al. 2011b). The application of the chase on $\mathcal{F}$ (i.e. $\text{C1}_{\mathcal{R}}(\mathcal{F})$) produces a saturated set of facts $\mathcal{F}^*$. We say a query is entailed from $\mathcal{K}$ iff $\text{C1}_{\mathcal{R}}(\mathcal{F}) \models Q$ where $\models$ is the FOL entailment of classical logic. Since the chase always halts, the set $\text{C1}_{\mathcal{R}}(F)$ for any $F \subseteq \mathcal{F}$ is finite, hence query entailment is **decidable**.

Inconsistency-Tolerant Semantics Given a knowledge base $\mathcal{K} = (\mathcal{F}, \mathcal{R}, \mathcal{N})$, a set $F \subseteq \mathcal{F}$ is said to be *inconsistent* iff $\text{C1}_{\mathcal{R}}(\mathcal{F}) \models \perp$, we say $\mathcal{K}$ is inconsistent iff $\mathcal{F}$ is inconsistent ($\mathcal{R}$ and $\mathcal{N}$ are assumed to be consistent). In presence of inconsistency every query can be entailed from $\mathcal{K}$. A common solution (Lembo et al. 2010) is to construct maximal (w.r.t $\subseteq$) consistent subsets of $\mathcal{F}$ called *repairs*, denoted by $\text{Repair}(\mathcal{K})$. Querying the repairs with different strategies yield different semantics.

Definition 1 (Lembo et al. 2010). Let $\mathcal{K} = (\mathcal{F}, \mathcal{R}, \mathcal{N})$ be a knowledge base and let Q be a query.

Brave : $\mathcal{K} \models_{\text{Brave}} Q$ iff there exists at least one repair $\mathcal{A}$ such that $\text{C1}_{\mathcal{R}}(\mathcal{A}) \models Q$.

IAAR : $\mathcal{K} \models_{\text{IAAR}} Q$ iff $\text{C1}_{\mathcal{R}}(\mathcal{I}) \models Q$ such that $\mathcal{I}$ is the intersection of all repairs.

Note that all semantics collapse to the classical entailment when the knowledge base is consistent. Note also that $\text{Repair}(\mathcal{K})$ is finite because $\mathcal{F}$ is finite and for all repairs $\mathcal{A}$, $\text{C1}_{\mathcal{R}}(\mathcal{A})$ is finite (it follows from the termination of the chase procedure (Baget et al. 2011a)).

Logic-based Argumentation

Now we shift to the definition of argumentation frameworks within our language. Let us first define an *argument*.

Definition 2 (Argument) Given (a possibly) inconsistent knowledge base $\mathcal{K}$. An *argument a* for a fact $\mathcal{C}$ w.r.t. $\mathcal{K}$ is a tuple $(\mathcal{H}, \mathcal{C})$ where $\mathcal{H}$ is a minimal consistent subset $\mathcal{H} \subseteq \mathcal{F}$ such that $\text{C1}_{\mathcal{R}}(\mathcal{H}) \models \mathcal{C}$.

For an argument $a = (\mathcal{H}, \mathcal{C})$, we denote by $\text{Hyp}(a) = \mathcal{H}$ and $\text{Conc}(a) = \mathcal{C}$ the hypothesis and the conclusion of a respectively. If $\text{Conc}(a) \models Q$ we say the argument a supports the query Q . For a set $F \subseteq \mathcal{F}$, $\text{Arg}(F) = \{a \mid a \text{ is an argument and } \text{Hyp}(a) \subseteq F\}$, hence $\text{Arg}(\mathcal{F})$ is the set of all arguments that can be constructed from the set of facts $\mathcal{F}$.

An attack between two arguments is defined as a binary relation $\mathcal{U}$ over the set of arguments $\text{Arg}(\mathcal{F})$, i.e. $a \mathcal{U} b$ means a attacks b . Attack captures inconsistency between the conclusion and hypothesis of arguments. This type of attack is called *undercut* or *assumption attack*.

Definition 3 (Attack) Let a and b be two distinct arguments. a attacks b ($a \mathcal{U} b$) iff $\exists H \subset \text{Hyp}(b)$ such that $\text{C1}_{\mathcal{R}}(\{\text{Conc}(a), H\}) \models \perp$.

Note that $\mathcal{U}$ is asymmetric (Croitoru and Vesic 2013) and does not contain self-attack.

We consider in this paper the class of finite argumentation frameworks because most of the results in the community concern this class (Amgoud, Besnard, and Vesic 2014).

Definition 4 (Argumentation framework) Given an inconsistent knowledge base $\mathcal{K}$. The corresponding argumentation framework of $\mathcal{K}$ is the tuple $\mathcal{S} = (\mathcal{G}, \mathcal{U})$ such that $\mathcal{G} = \text{Arg}(\mathcal{F})$ and $\mathcal{U}$ is the set of attacks between all arguments in $\mathcal{G}$.

The argumentation framework $\mathcal{S} = (\mathcal{G}, \mathcal{U})$ is finite. This follows from the fact that $\forall a \in \mathcal{G}, \text{Hyp}(a) \subseteq \mathcal{F}$ where $\mathcal{F}$ is finite. Note that Hence $\mathcal{G}$ has a bounded set of arguments.

An important relation between *IAAR* and *Brave* semantics and the so-called causes has been shown in (Bienvenu, Bourgaux, and Goasdoué 2016). It turns out that causes coincide with the definition of arguments. In fact, this relation is very important to establish the results in Section 4.

Proposition 1 Let $\mathcal{K}$ be an inconsistent knowledge base and $\mathcal{S} = (\mathcal{G}, \mathcal{U})$ its corresponding argumentation framework and Q a query. The following holds:

- $\mathcal{K} \models_{\text{Brave}} Q$ iff $\exists a \in \mathcal{G}$ such that $\text{Conc}(a) \models Q$.
- $\mathcal{K} \models_{\text{IAAR}} Q$ iff $\exists a \in \mathcal{G}$ such that $\text{Conc}(a) \models Q$ and $\nexists b \in \mathcal{G}, (b, a) \in \mathcal{U}$.

In our work we use argumentation as a formal framework to study and characterize the *IAAR* and *Brave* semantics in terms of argumentation dialogue between participants with different profiles. In what follows we detail this contribution.

Argumentation Dialogue System $\mathcal{A}_0$

Our aim in this section is to define a **minimal** (in terms of the dialogue's protocol) formalization of argumentation dialogues (denoted $\mathcal{A}_0$) of (Amgoud, Saint-Cyr, and Dupin 2013). Given an inconsistent knowledge base $\mathcal{K}$, $\mathcal{A}_0$ is a turn-taking dialogue game between two participants *proponent* (PRO) and *opponent* (OPP) arguing in favor or against a query in $\mathcal{K}$ (the subject or the thesis). They have both direct access to the same knowledge base $\mathcal{K}$. The participants exchange only arguments and counterarguments (called moves). The turn in this dialogue shifts in a non-deterministic way, precisely when one of the participant makes her/his point (Prakken 2006). In what follows we define what is an argumentation dialogue in the general sense. Then we show the restricted definition within the argumentation dialogue system $\mathcal{A}_0$.

Definition 5 (Dialogues) A dialogue $\mathcal{D}_n = (a_0, a_1, \dots, a_n)$ is a sequence of arguments such that

$\text{Player}(a_0) = \text{PRO}$ where $\text{Player}(a_i)$ for all $a_i \in \mathcal{D}_n$ denotes the participant who plays the argument a_i . For all $a_i \in \mathcal{D}_n$, $i > 0$, $\text{Player}(a_i)$ is either PRO or OPP. a_n is called the most recent argument in $\mathcal{D}_n$. The subject of the dialogue is a query Q and it is denoted as $\text{Subject}(\mathcal{D}_n)$. We denote by $\mathcal{D}$ the set of all dialogues that can be generated over a given $\mathcal{K}$.

This definition presents minimal requirements for a normal argumentation dialogue. In order to introduce argumentation dialogues of our system A_0 we impose *appropriateness* and *meaningful dispute*.

Definition 6 (Appropriateness) Given a dialogue $\mathcal{D}_n$. $\mathcal{D}_n$ is appropriate if and only if $\forall a_i \in \mathcal{D}_n$, either $\text{Conc}(a_i) \models \text{Subject}(\mathcal{D}_n)$ or $\exists a_j \in \mathcal{D}_n$, $0 \leq j < i$ such that a_i attacks a_j . a_j is denoted as $\text{Target}(a_i)$.

We call the first argument a support move and we call the second argument an attack move. If $\text{Target}(a_i)$ is a support move then we say a_i disqualifies $\text{Target}(a_i)$.

Appropriateness means that any advanced argument in the dialogue should either attack a previous argument or support the subject of the dialogue.

Definition 7 (Meaningful dispute) Given $\mathcal{D}_n$. $\mathcal{D}_n$ is meaningful if and only if for all support moves $a_j \in \mathcal{D}_n$, there exists no support move $a_i \in \mathcal{D}_n$, $0 \leq j < i$ such that $a_j = a_i$ and a_i is disqualified.

Meaningful dispute dictates that if a_i is a support move and it is disqualified then it is forbidden to be reused in further exchanges.

Definition 8 (Dialogues of A_0) Given an argumentation dialogue $\mathcal{D}_n$, we say $\mathcal{D}_n$ is an argumentation dialogue of A_0 if and only if it is appropriate and meaningful.

When one of the participants plays a move a_i we distinguish the following replies:

SUBJECT-SUPP : if $\text{Target}(a_i)$ is empty.

ATTACK-SUPPORT: if $\text{Target}(a_i)$ is a support move.

ATTACK-ATTACK: if $\text{Target}(a_i)$ is an attack move.

The first reply is called *subject-support* because the advanced move supports the subject of the dialogue and does not attack any arguments. The second reply is called *attack-support* because it attacks a support move. The last reply is called *attack-attack* reply because it attacks an attack move.

From now on when we use the general term *dialogue* we refer to those dialogues that respect previous definitions (i.e. dialogues of A_0).

Let us define when does any arbitrary dialogue terminate.

Definition 9 (Termination rule) A dialogue $\mathcal{D}_n$ terminates when neither PRO nor OPP can advance an argument.

Finiteness is a desirable property in formal argumentation dialogues (Amgoud, Saint-Cyr, and Dupin 2013), it ensures that termination is always guaranteed. In our definition of argumentation dialogue finiteness strongly depends on the participant's profile. A profile is briefly the way the participant prefers to reply.

Definition 10 (PRO's profiles) PRO's profile fits the following categories:

- (1) We say PRO has an **aggressive** profile if he always prefers ATTACK-ATTACK replies.
- (2) We say PRO has a **focused** profile if:
 - (a) he plays SUBJECT-SUPP replies (no ATTACK-ATTACK replies even if they are available); and,
 - (b) he plays only when all previously advanced support moves are disqualified.

An aggressive proponent tries first to counterattack the opponent's attack moves before advancing any support move. A focused proponent tries always to support the subject of the dialogue by advancing **only** support moves. He does so if and only if the previous support moves have been disqualified.

Definition 11 (OPP's profiles) OPP's profile fits the following categories:

- (1) We say OPP has an **aggressive** profile if she has no preferences over ATTACK-ATTACK or ATTACK-SUPPORT replies.
- (2) We say OPP has a **focused** profile if she plays only ATTACK-SUPPORT replies.

An aggressive opponent tries always to oppose the proponent just for the sake of argument. She consequently attacks without distinction between support moves or attack moves of PRO. While a focused opponent her main goal is to disqualify all proponent's support moves.

Let us study termination in the light of these profiles.

Proposition 2 Let $\mathcal{D}$ be all the dialogues over an arbitrary inconsistent knowledge $\mathcal{K}$. Then, $\forall \mathcal{D} \in \mathcal{D}$:

- (1) $\mathcal{D}$ terminates if at least one of the participants is **focused**.
- (2) $\mathcal{D}$ may not terminate if the two participants are **aggressive**.

Proof

- (1): let us suppose that PRO is focused. That means that he plays only SUBJECT-SUPP replies when previous support moves have been disqualified. Since the argumentation framework is finite, the set of support moves is finite. Hence, PRO eventually will stop using SUBJECT-SUPP replies (he runs out of support moves or the other cannot respond) thus the dialogue terminates. Let us suppose that OPP is focused. This means that OPP advances only ATTACK-SUPPORT replies. Since the set of support moves is finite and the set of attackers of any argument is finite then OPP will advance a finite ATTACK-SUPPORT replies. When OPP runs out of ATTACK-SUPPORT replies the dialogue terminates.
- (2): suppose that PRO and OPP are aggressive, that means PRO advances always ATTACK-ATTACK replies and PRO may advance ATTACK-ATTACK or ATTACK-SUPPORT replies. Let us suppose that in the worst case OPP advances only ATTACK-ATTACK. In this case, consider $\mathcal{AF}_{\mathcal{K}} = (\{a, b\}, \{(a, b), (b, a)\})$. It is clear that the two

parties will get into an infinite loop of attack and counterattack. PRO advances a then OPP replies with b and so on and so forth. Thus the dialogue is not guaranteed to terminate.

■

The reason that the dialogue may not terminate when the two participants are aggressive is that they may get into a loop of arguments and counterarguments without focusing on the subject of the dialogue. The intuition behind termination is that if they are focused then they will advance a finite set of moves and they avoid to get into possibly infinite loops.

Let us decide when does a participant win the dialogue. Note that the winner determination in A_0 is also minimal. More sophisticated conditions, which are not necessary in our context, can be found in (Prakken 2006; Walton and Krabbe 1995)).

Definition 12 (Winner) Let $\mathcal{D}_n$ be a terminated dialogue. PRO wins $\mathcal{D}_n$ if and only if there exists a support move $a_i \in \mathcal{D}_n$ that is not disqualified by OPP. Otherwise OPP wins.

The winning criterion dictates that PRO wins if he successfully provided an argument that supports the subject of $\mathcal{D}_n$.

Relation between A_0 and Inconsistency-tolerant Semantics

In this section we look at the relation between the semantics *Brave* and *JAR* with the argumentation dialogues of A_0 (Section 3). We look first on what would be the outcome of the dialogue (the winner) given a query that conforms a given semantics. We look next on what would be the semantics a given query conforms if it is the subject of a given dialogue with an arbitrary outcome.

The following evident but yet important observation can be stated.

Observation 1 (Consistent winner) Let $\mathcal{K}$ be a consistent knowledge base and Q a query. if $\mathcal{K} \models Q$ then any dialogue $\mathcal{D}_n$ over $\mathcal{K}$ with subject Q terminates in one step with PRO as the winner.

Note that we do not consider the case where $\mathcal{K} \not\models Q$ because in this case the dialogue won't even start.

Let us see who would be the winner if the query is *Brave*-entailed and not *JAR*-entailed.

Proposition 3 (Brave winner) Let $\mathcal{K}$ be an inconsistent knowledge base such that $\mathcal{K} \models_{\text{Brave}} Q$ and $\mathcal{K} \not\models_{\text{JAR}} Q$. Then OPP wins any dialogue $\mathcal{D}_n$ over $\mathcal{K}$ with subject Q if:

- (1) OPP is focused and PRO is either focused or aggressive.
- (2) OPP is aggressive and PRO is focused.
- (3) OPP and PRO are both aggressive and $\mathcal{D}_n$ terminates.

Proof

- (1): if $\mathcal{K} \models_{\text{Brave}} Q$ then there exists an argument a that supports Q (Prop 1). Furthermore, if $\mathcal{K} \not\models_{\text{JAR}} Q$ then there exists at least one repair $\mathcal{A}$ such that $a \notin \text{Arg}(\mathcal{A})$ where $\text{Arg}(\mathcal{A})$ is the set of all arguments that can be

constructed from $\mathcal{A}$. If $a \notin \text{Arg}(\mathcal{A})$ then there exists an argument b in $\text{Arg}(\mathcal{A})$ such that $b \mathcal{U} a$ (by maximality of repairs). Hence, whenever PRO plays a support move for Q then OPP will always be able to reply by an ATTACK-SUPPORT reply. Since OPP is focused (plays only ATTACK-SUPPORT) eventually all PRO's support moves will be disqualified. The dialogue is guaranteed to terminate (Prop 2), so OPP wins eventually.

- (2): in this case $\mathcal{D}_n$ terminates (Prop 2). Given $\mathcal{K} \models_{\text{Brave}} Q$ and $\mathcal{K} \not\models_{\text{JAR}} Q$ then either there is a support move which is not disqualified or all the support moves are disqualified. It follows from (1) that the former is impossible, therefore all the support moves are disqualified, thus OPP wins.
- (3): it follows from (2).

■

If the query is *JAR*-entailed then the winner changes as follows.

Proposition 4 (JAR winner) Let $\mathcal{K}$ be an inconsistent knowledge base such that $\mathcal{K} \models_{\text{JAR}} Q$. Then PRO wins any dialogue $\mathcal{D}_n$ over $\mathcal{K}$ with subject Q if:

- (1) PRO is focused and OPP is either focused or aggressive.
- (2) PRO is aggressive and OPP is focused.
- (3) PRO and OPP are both aggressive and $\mathcal{D}_n$ terminates.

Proof

- (1): it follows from Proposition 1. If $\mathcal{K} \models_{\text{JAR}} Q$ then there is an argument a that supports Q for which there is no $b \in \mathcal{G}, b \mathcal{U} a$. Therefore, PRO will eventually play the support move a against which OPP has no attack. Hence, the dialogue terminates and PRO wins. (2) and (3) follows from (1) because $\mathcal{D}_n$ terminates under these conditions.

■

One can clearly see that the profile “focused” corresponds to a dominant winning strategy in game theory with respect to *Brave* and *JAR* semantics. PRO wins always in *JAR* if he is focused and OPP wins always in *Brave* if he is focused.

Now let us take the inverse direction where we define the semantics that conforms a given dialogue. We show that *JAR* and the winning of the dialogue are equivalent. Note that we assume that they do not hide arguments.

Proposition 5 (JAR equivalence) Given a $\mathcal{D}_n$ with subject Q over an inconsistent $\mathcal{K}$. PRO wins $\mathcal{D}_n$ iff $\mathcal{K} \models_{\text{JAR}} Q$.

Here whatever the profile of the participants, if the dialogue terminates and PRO wins this means that there is a support move which is not disqualified (not attacked). In this case we have found an argument a which is not attacked. This is exactly the requirement defined in Proposition 1.

Conclusion

This paper provides a dialectical characterization of *JAR* and *Brave* semantics. This was done via the A_0 argumentation dialogue system considering different participants profiles. While the state of the art never considered these semantics from a dialectical point of view, an alternative way of

characterizing these semantics would be via the dialectical proof approaches of (Modgil and Caminada 2009). Knowing that $\mathcal{IAR}$ corresponds to the grounded semantics (Croitoru and Vesic 2013) means that such approaches are a good candidate for investigation. However the reason why we did not take this route lies in the way they characterize the semantics (Modgil and Caminada 2009). They do so by varying the rules of the dialogue. We instead consider participants' profiles while keeping the rules invariable. In addition their work is applied on abstract argumentation framework while we work with logic-based argumentation framework. Having said this, instantiating their work remains a possible alternative manner to consider, especially if the semantics are $\mathcal{AR}$ and $\mathcal{ICR}$. In this case equivalence results with the results presented in this paper might be insightful.

Acknowledgments

The authors acknowledge the support of ANR grant DUR-DUR (ANR-13-ALID-0002). The work of the second author has been carried out in a part of the research delegation at INRA MISTEA Montpellier and INRA IATE CEPIA Axe 5 Montpellier.

References

- Amgoud, L.; Besnard, P.; and Vesic, S. 2014. Equivalence in logic-based argumentation. *Journal of Applied Non-Classical Logics* 24(3):181–208.
- Amgoud, L.; Saint-Cyr, D.; and Dupin, F. 2013. An axiomatic approach for persuasion dialogs. In *Tools with Artificial Intelligence (ICTAI), 2013 IEEE 25th International Conference on*, 618–625. IEEE.
- Arioua, A.; Tamani, N.; Croitoru, M.; and Buche, P. 2014. Query failure explanation in inconsistent knowledge bases: A dialogical approach. In *Research and Development in Intelligent Systems XXXI*. Springer. 119–133.
- Arioua, A.; Tamani, N.; and Croitoru, M. 2015. Query answering explanation in inconsistent datalog +/- knowledge bases. In *Database and Expert Systems Applications - 26th International Conference, DEXA 2015*, 203–219.
- Baget, J.-F.; Mugnier, M.-L.; Rudolph, S.; and Thomazo, M. 2011a. Walking the complexity lines for generalized guarded existential rules. In *Proc. of IJCAI'11*, 712–717.
- Baget, J.; Leclère, M.; Mugnier, M.; and Salvat, E. 2011b. On rules with existential variables: Walking the decidability line. *Artif. Intell.* 175(9-10):1620–1654.
- Bienvenu, M.; Bourgaux, C.; and Goasdoué, F. 2016. Explaining query answers under inconsistency-tolerant semantics over description logic knowledge bases. In *Proc of AAAI'16 (to appear)*.
- Bienvenu, M. 2012. On the complexity of consistent query answering in the presence of simple ontologies. In *Proc of AAAI'12*, 705–711.
- Croitoru, M., and Vesic, S. 2013. What can argumentation do for inconsistent ontology query answering? In *Scalable Uncertainty Management*, volume 8078 of *Lecture Notes in Computer Science*, 15–29. Springer Berlin Heidelberg.
- Lembo, D.; Lenzerini, M.; Rosati, R.; Ruzzi, M.; and Savo, D. F. 2010. In *proc of RR'10*, 103–117. Springer-Verlag.
- Modgil, S., and Caminada, M. 2009. Proof theories and algorithms for abstract argumentation frameworks. In *Argumentation in artificial intelligence*. Springer. 105–129.
- Parsons, S.; Wooldridge, M.; and Amgoud, L. 2003. On the outcomes of formal inter-agent dialogues. In *Proceedings of the second international joint conference on Autonomous agents and multiagent systems*, 616–623. ACM.
- Prakken, H. 2006. Formal systems for persuasion dialogue. *The Knowledge Engineering Review* 21(02):163–188.
- Walton, D., and Krabbe, E. C. 1995. *Commitment in dialogue: Basic concepts of interpersonal reasoning*. SUNY press.