



A Replication Study of Semantics in Argumentation

Leila Amgoud

► To cite this version:

Leila Amgoud. A Replication Study of Semantics in Argumentation. Twenty-Eighth International Joint Conference on Artificial Intelligence (IJCAI 2019), Aug 2019, Macao, China. pp.6260-6266, <10.24963/ijcai.2019/874>. <hal-02325890>

HAL Id: hal-02325890

<https://hal.science/hal-02325890v1>

Submitted on 22 Oct 2019

HAL is a multi-disciplinary open access archive for the deposit and dissemination of scientific research documents, whether they are published or not. The documents may come from teaching and research institutions in France or abroad, or from public or private research centers.

L'archive ouverte pluridisciplinaire **HAL**, est destinée au dépôt et à la diffusion de documents scientifiques de niveau recherche, publiés ou non, émanant des établissements d'enseignement et de recherche français ou étrangers, des laboratoires publics ou privés.



HAL Authorization

A Replication Study of Semantics in Argumentation

Leila Amgoud

CNRS – IRIT, France

Leila.Amgoud@irit.fr

Abstract

Argumentation aims at increasing acceptability of claims by supporting them with *arguments*. Roughly speaking, an argument is a set of premises intended to establish a definite claim. Its strength depends on the plausibility of the premises, the nature of the link between the premises and claim, and the prior acceptability of the claim. It may generally be weakened by other arguments that undermine one or more of its three components.

Evaluation of arguments is a crucial task, and a sizable amount of methods, called *semantics*, has been proposed in the literature. This paper discusses two classifications of the existing semantics: the first one is based on the type of semantics' outcomes (sets of arguments, weighting, and preorder), the second is based on the goals pursued by the semantics (acceptability, strength, coalitions).

1 Introduction

Argumentation is a reasoning approach based on the justification of claims by *arguments*, i.e. reasons for accepting claims. It received great interest from the Artificial Intelligence community since late 1980s, namely as a unifying approach for nonmonotonic reasoning [Lin and Shoham, 1989]. It was later used for solving different other problems like reasoning with inconsistent information (eg. [Simari and Loui, 1992; Besnard and Hunter, 2001]), decision making (eg. [Zhong *et al.*, 2019]), classification (eg. [Amgoud and Serrurier, 2008]), etc. It has also several practical applications, namely in legal and medical domains (see [Atkinson *et al.*, 2017]).

Whatever the problem to be solved, an argumentation process follows generally four main steps: to justify claims by arguments, identify (attack, support) relations between arguments, evaluate the arguments, and define an output. The last step depends on the results of the evaluation. For instance, an inference system draws formulas that are justified by what is qualified at the evaluation step as “strong” arguments. Evaluation of arguments is thus crucial as it impacts the outcomes of argument-based systems. Consequently, a plethora of methods, called *semantics*, have been proposed in the literature. The very first ones are the *extension* semantics (*stable*,

preferred, *complete* and *grounded*) that were proposed in the seminal paper [Dung, 1995]. These semantics were refined by several scholars (eg. *recursive* [Baroni *et al.*, 2005], *ideal* [Dung *et al.*, 2007], *semi-stable* [Caminada, 2006]). In [Cayrol and Lagasque-Schiex, 2005] another type of semantics, called *gradual* or *weighting*, was introduced with the purpose of refining the above cited semantics. Examples of such semantics are Trust-based [da Costa Pereira *et al.*, 2011], social semantics [Leite and Martins, 2011], Iterative Schema [Gabbay and Rodrigues, 2015], Categorizer [Pu *et al.*, 2014], (DF)-QuAD [Baroni *et al.*, 2015; Rago *et al.*, 2016], Max-based, Card-based and Weighted Categorizer [Amgoud *et al.*, 2017]. Finally, *ranking* semantics were defined in [Amgoud and Ben-Naim, 2013] and examples of such semantics are Burden-based and Discussion-based semantics [Amgoud and Ben-Naim, 2013], the propagation-based semantics [Bonzon *et al.*, 2016; 2017] and the ones from [Dondio, 2018]. The above semantics were analyzed and compared on the basis of postulates proposed in [Amgoud and Ben-Naim, 2016a; 2016b]. The results show that they follow different principles. Extension semantics from [Dung, 1995] were also empirically analyzed in [Rahwan *et al.*, 2010; Rosenfeld and Kraus, 2014; 2015]. These studies revealed that humans do not follow the main principles behind those semantics, namely they argued against the cognitive plausibility of *reinstatement* principle. Another empirical study from [Rosenfeld and Kraus, 2016] argued in favor of weighting semantics.

Focusing on attack argumentation graphs, this paper presents a survey of existing semantics. It proposes two classifications of those semantics. The first one is based on the type of semantics' outcomes. It shows that there are three families of semantics: extension semantics that return sets of arguments, weighting semantics that assign a single value to every argument, and ranking semantics that return a total or partial preorder on the set of arguments. The second classification clarifies the nature itself of the outcome, or the goal pursued by a semantics. We argue that weighting and ranking semantics evaluate the strength of individual arguments in graphs, while extension semantics look for acceptable arguments, i.e., those that a rational agent can accept.

2 Classification wrt Type of Outcome

An *argumentation graph*, called also *argumentation framework* in [Dung, 1995], is a set of arguments and a binary re-

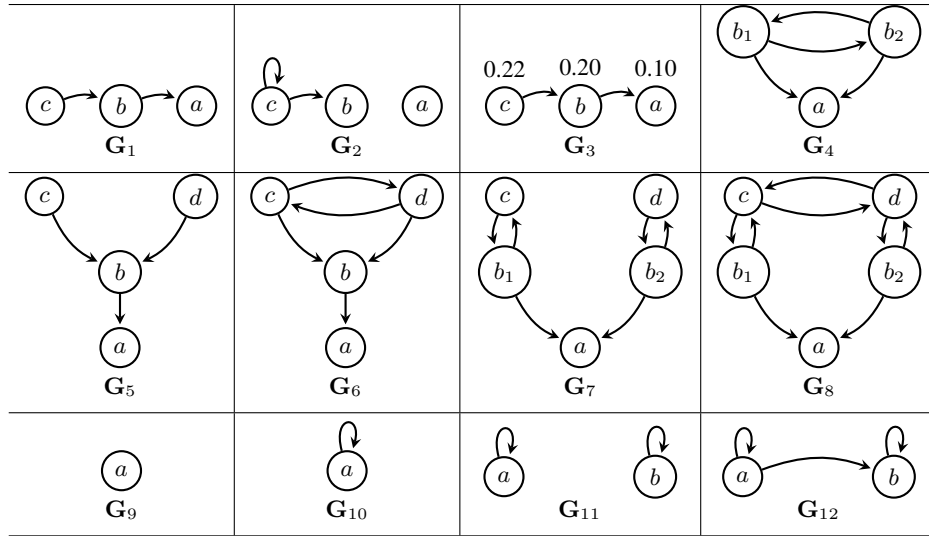


Table 1: Examples of argumentation graphs

lation representing attacks between the arguments. An attack expresses that an argument supports the fact that one of the components (premises, claim, link) of the targeted argument does not hold. In what follows, arguments are considered abstract entities whose internal structure is not specified. Let Arg denote the universe of all possible arguments.

Definition 1 An argumentation graph is an ordered pair $\mathbf{G} = \langle \mathcal{A}, \mathcal{R} \rangle$, where $\mathcal{A} \subseteq \text{Arg}$ is a non-empty and finite set and $\mathcal{R} \subseteq \mathcal{A} \times \mathcal{A}$. For $a, b \in \mathcal{A}$, $(a, b) \in \mathcal{R}$ means: a attacks b . Let AG be the set of all argumentation graphs built from Arg .

The plethora of semantics that have been proposed in the computational argumentation literature can be organized into three classes according to the *type of their outcomes*.

- *Extension semantics* return subsets of arguments.
- *Weighting semantics* ascribe a (numerical or qualitative) value to every argument.
- *Ranking semantics* return a preorder (i.e., a reflexive and transitive relation) on the set of arguments.

Definition 2 A semantics is a function π that assigns to every argumentation graph $\mathbf{G} = \langle \mathcal{A}, \mathcal{R} \rangle \in \text{AG}$,

- a set $\text{Ext}_{\mathbf{G}}^{\pi} \subseteq \mathcal{P}(\mathcal{A})$ (Extension Semantics)
- a weighting $\text{Deg}_{\mathbf{G}}^{\pi} : \mathcal{A} \rightarrow \mathcal{D}$ (Weighting Semantics)
- a preorder $\succeq_{\mathbf{G}}^{\pi} \subseteq \mathcal{A} \times \mathcal{A}$ (Ranking Semantics)

where $\mathcal{P}(\mathcal{A})$ stands for the powerset of $\mathcal{A}$ and $\mathcal{D}$ is a totally ordered scale.

2.1 Extension Semantics

Initiated in [Dung, 1995], extension semantics identify arguments that are *acceptable* for (thus can be accepted by) a rational agent. The following particular definition was used:

An argument is acceptable for a rational agent if it can be defended against all attacks on it. (1)

Dung proposed different ways of defining formally this notion of acceptability. They are all based on the same idea:

identifying sets of arguments, called *extensions*, that defend their elements against all attacks. Each extension represents an *alternative set of acceptable arguments*. Extension semantics are grounded on three crucial notions: conflict-freeness, defence, and extensions.

Definition 3 Let $\mathbf{G} = \langle \mathcal{A}, \mathcal{R} \rangle \in \text{AG}$ and $\mathcal{E} \subseteq \mathcal{A}$.

- $\mathcal{E}$ is conflict-free iff $\nexists a, b \in \mathcal{E}$ such that $(a, b) \in \mathcal{R}$.
- $\mathcal{E}$ defends $a \in \mathcal{A}$ iff $\forall b \in \mathcal{A}$, if $(b, a) \in \mathcal{R}$, then $\exists c \in \mathcal{E}$ such that $(c, b) \in \mathcal{R}$.

We recall below the extension semantics proposed in [Dung, 1995]. Interested readers can refer to [van der Torre and Vesic, 2017] for a complete presentation of all existing extension semantics as well as their formal analysis.

Definition 4 Let $\mathbf{G} = \langle \mathcal{A}, \mathcal{R} \rangle \in \text{AG}$, $\mathcal{E} \subseteq \mathcal{A}$ is conflict-free.

- $\mathcal{E}$ is a complete extension iff it defends all its elements and contains any argument it defends.
- $\mathcal{E}$ is a preferred extension iff it is a maximal (w.r.t. set inclusion) complete extension.
- $\mathcal{E}$ is a stable extension iff it attacks any argument in $\mathcal{A} \setminus \mathcal{E}$.
- $\mathcal{E}$ is a grounded extension iff it is a minimal (w.r.t. set inclusion) complete extension.

Example 1 Under stable semantics, the graph $\mathbf{G}_1$ of Table 1 has one extension $\{a, c\}$. Both a, c are acceptable while b is unacceptable. $\mathbf{G}_2$ has no extension, then no evaluation can be done. Note that a is acceptable under preferred semantics. $\mathbf{G}_4$ has two sets of acceptable arguments: $\{b_1\}$ and $\{b_2\}$.

Conditions for Acceptability

According to the above semantics, for an argument to be acceptable in the sense (1), it should satisfy two conditions:

- *Being defended:* Consider the graphs $\mathbf{G}_4$ and $\mathbf{G}_7$. Under stable semantics, denoted here by s , $\text{Ext}_{\mathbf{G}_4}^s = \{\mathcal{E}_1, \mathcal{E}_2\}$ with $\mathcal{E}_1 = \{b_1\}$ and $\mathcal{E}_2 = \{b_2\}$, and $\text{Ext}_{\mathbf{G}_7}^s = \{\mathcal{E}'_1, \mathcal{E}'_2, \mathcal{E}'_3, \mathcal{E}'_4\}$,

$f_{comp}(x_1, x_2) = x_1(1 - x_2)$	$g_{sum}(x_1, \dots, x_n) = \sum_{i=1}^n x_i$
$f_{exp}(x_1, x_2) = x_1 e^{-x_2}$	$g_{sum, \alpha}(x_1, \dots, x_n) = \left(\sum_{i=1}^n (x_i)^\alpha \right)^{\frac{1}{\alpha}}$
$f_{frac}(x_1, x_2) = \frac{x_1}{1+x_2}$	$g_{max}(x_1, \dots, x_n) = \max\{x_1, \dots, x_n\}$

Table 2: Examples of functions f and g

with $\mathcal{E}'_1 = \{c, d, a\}$, $\mathcal{E}'_2 = \{c, b_2\}$, $\mathcal{E}'_3 = \{d, b_1\}$, $\mathcal{E}'_4 = \{b_1, b_2\}$. In both graphs, the two attackers of a are attacked and defend themselves against those attacks. Thus, they are acceptable in that they belong to some extensions. However, a is unacceptable in $\mathbf{G}_4$ while it is acceptable in $\mathbf{G}_7$ since it belongs to $\mathcal{E}'_1$. The main difference between the two graphs is the fact that a is defended in $\mathbf{G}_7$ and not defended in $\mathbf{G}_4$.

- *Belonging to an extension:* The fact of being defended against all attacks is not sufficient for being acceptable. An argument should in addition belong to an extension. The graph $\mathbf{G}_8$ has three stable extensions: $\mathcal{E}''_1 = \{b_1, b_2\}$, $\mathcal{E}''_2 = \{b_1, d\}$, $\mathcal{E}''_3 = \{b_2, c\}$. Unlike in graph $\mathbf{G}_7$, the argument a is unacceptable in $\mathbf{G}_8$ even if it is defended against its two attackers exactly as in $\mathbf{G}_7$. It is worth mentioning that the rejection of a in $\mathbf{G}_8$ is not due to the fact that its two defenders (c and d) are in conflict. In $\mathbf{G}_6$ the two defenders of a are conflicting but a still belongs to the two stable extensions of the graph, and is thus acceptable.

Specificities of the Approach

The extension-based approach for semantics has two particularities. First, the acceptability of an argument belonging to an argumentation graph depends on the *topology of the whole graph* and not on the acceptability of the direct attackers of the argument. Indeed, an argument may be rejected even if all its direct attackers are themselves rejected. This is, for example, the case of arguments of any odd-length cycle under preferred semantics. Second, the approach provides a *complete view* on *all* arguments of a graph at once. Indeed, it identifies the set(s) of all acceptable arguments.

Gradual Acceptability

Recall that the purpose of extension semantics is to answer the question on whether a given argument is acceptable (i.e., it can be accepted by a rational agent). While grounded semantics returns a single extension, thus an argument is acceptable if it belongs to it and unacceptable otherwise, the three other semantics may lead to multiple extensions. Hence, the status of an argument regarding acceptability becomes less clear. Consequently, different *levels* of acceptability were introduced in the literature for facilitating the exploitation of extension semantics in concrete applications. For instance in [Cayrol and Lagasque-Schiex, 2005], an argument is:

- skeptically accepted if it belongs to all extensions.
- cleanly accepted if it belongs to some extension and none of its attackers belong to an extension.
- credulously accepted if it belongs to some extension.
- rejected if it does not belong to any extension.

An argument is acceptable if it is either skeptically or cleanly accepted, it is unacceptable otherwise. Skeptically

accepted arguments are considered as more acceptable than cleanly accepted ones. Credulously accepted arguments are unacceptable since their direct attackers, which point out some flaws in the arguments, can themselves be credulously accepted. Taking advantage of the labeling approach of extension semantics, [Bonzon *et al.*, 2018] proposed five levels of acceptability depending on the labels (in, out, undec) assigned to each argument. We present them from the strongest level to the weakest one: $\{\text{in}\} \succ \{\text{in}, \text{undec}\} \succ \{\text{undec}\} \approx \{\text{in}, \text{out}, \text{undec}\} \approx \{\text{in}, \text{out}\} \succ \{\text{out}, \text{undec}\} \succ \{\text{out}\}$.

2.2 Weighting Semantics

Introduced for the first time in [Cayrol and Lagasque-Schiex, 2005] under the name of *gradual valuation methods*, weighting semantics assign a value from a fixed ordered scale to every argument in a graph. The value represents the *strength* of the argument. For any argumentation graph $\mathbf{G} = \langle \mathcal{A}, \mathcal{R} \rangle$, any argument $a \in \mathcal{A}$, the value (or strength) of a under a given weighting semantics π is defined as follows:

$$\text{Deg}_{\mathbf{G}}^{\pi}(a) = f(g(\text{Deg}_{\mathbf{G}}^{\pi}(b_1), \dots, \text{Deg}_{\mathbf{G}}^{\pi}(b_n))),$$

where $b_1, \dots, b_n$ are the direct attackers of a , g is an *aggregation* function that evaluates how strongly a is attacked, and f is an *influence* function that takes into account the strength of attacks. Examples of functions studied in the literature are given in Table 2 (see [Amgoud and Doder, 2019] for more functions). We recall in the next definition some weighting semantics (h -Categorizer from [Besnard and Hunter, 2001], Max-based [Amgoud *et al.*, 2017], and compensation semantics from [Amgoud *et al.*, 2016]). The two first semantics use the scale $\mathcal{D} = [0, 1]$ while the third one uses $\mathcal{D} = [1, +\infty)$.

Definition 5 Let $\mathbf{G} = \langle \mathcal{A}, \mathcal{R} \rangle \in \text{AG}$ and $a \in \mathcal{A}$. $\text{Deg}_{\mathbf{G}}^{\pi}(a) =$

$$\begin{cases} \frac{1}{1 + \sum_{(b,a) \in \mathcal{R}} \frac{1}{\text{Deg}_{\mathbf{G}}^{\pi}(b)}} & (h\text{-Categorizer}) \\ \frac{1}{1 + \max_{(b,a) \in \mathcal{R}} \frac{1}{\text{Deg}_{\mathbf{G}}^{\pi}(b)}} & (\text{Max-based}) \\ 1 + \left(\sum_{(b,a) \in \mathcal{R}} \frac{1}{(\text{Deg}_{\mathbf{G}}^{\pi}(b))^{\alpha}} \right)^{1/\alpha}, \alpha \in (0, +\infty) & \end{cases}$$

h -Categorizer, Max-based, and compensation semantics use the same influence function f_{frac} . However, their aggregation functions are respectively g_{sum} , g_{max} and $g_{sum, \alpha}$.

Example 1 (Cont) Consider h -Categorizer. In graph $\mathbf{G}_1$, $\text{Deg}_{\mathbf{G}_1}^{\pi}(a) = 1$, $\text{Deg}_{\mathbf{G}_1}^{\pi}(b) = 0.5$ and $\text{Deg}_{\mathbf{G}_1}^{\pi}(c) = 0.66$. In $\mathbf{G}_2$, $\text{Deg}_{\mathbf{G}_2}^{\pi}(a) = 1$ and $\text{Deg}_{\mathbf{G}_2}^{\pi}(b) = \text{Deg}_{\mathbf{G}_2}^{\pi}(c) = 0.61$. Note that $\text{Deg}_{\mathbf{G}_5}^{\pi}(a) = 0.75$ while $\text{Deg}_{\mathbf{G}_6}^{\pi}(a) = 0.69$, this means that the strengths of a 's defenders are taken into account.

Unlike extension semantics, the existing weighting semantics focus on *direct attackers* of arguments and not on the whole topology of the graph; focus on *individual* arguments rather than sets of arguments; return a *single output* (a vector of values). However, they are silent on which arguments to accept. *h*-Categorizer ascribes one value for each of the three arguments in $\mathbf{G}_1$, but it does not specify which argument should be accepted and which one should be rejected. In a next section, we will explain the origin of these differences between extension semantics and weighting semantics.

2.3 Ranking Semantics

Ranking semantics were introduced in [Amgoud and Ben-Naim, 2013] with the aim of introducing *graduality in acceptability*, and thus to rank order arguments from the most to the least acceptable ones. The authors started by providing a list of properties (called principles) that a ranking semantics should satisfy, among which *Void Precedence* (VP) and *Counter-Transitivity* (CT). (VP) states that an argument that has no attackers is more acceptable than any attacked argument. (CT) states that an argument a should be at least as acceptable as an argument b if the attackers of b are at least as numerous and as acceptable as those of a . The authors proposed then two ranking semantics: Burden and Discussion based. Propagation semantics proposed in [Bonzon *et al.*, 2016; 2017] and the one from [Dondio, 2018] are other examples of ranking semantics. We recall below one semantics proposed in [Bonzon *et al.*, 2016]. Its basic idea is to give some *power* to non-attacked arguments by ascribing initial values to arguments. Before introducing the semantics, let us recall the definition of lexicographical order.

Let $V = \langle V_1, \dots, V_n \rangle$ and $V' = \langle V'_1, \dots, V'_m \rangle$ be two vectors of real numbers. $V >_{lex} V'$ iff $\exists i \leq n$ such that $V_i > V'_i$ and $\forall j < i, V_j = V'_j$. $V \geq_{lex} V'$ means it is not the case that $V' >_{lex} V$.

Definition 6 Let $\mathbf{G} = \langle \mathcal{A}, \mathcal{R} \rangle \in \mathbf{AG}$, $v : \mathcal{A} \rightarrow \{\epsilon, 1\}$ where $\epsilon \in [0, 1)$ and $\forall a \in \mathcal{A}, v(a) = 1$ if a is not attacked and $v(a) = \epsilon$ else. The value of $a \in \mathcal{A}$ at step $i \in \mathbb{N}$ is $P_i(a)$ s.t.:

$$P_i(a) = \begin{cases} v(a) & \text{if } i = 0 \\ P_{i-1}(a) + (-1)^i \sum_{(b,a) \in \mathcal{R}} v(b) & \text{else} \end{cases}$$

The propagation vector of a is $P(a) = \langle P_0(a), P_1(a), \dots \rangle$. For $a, b \in \mathcal{A}$, a is at least as acceptable as b , denoted by $a \succeq_{\mathbf{G}}^P b$, iff $P(a) \geq_{lex} P(b)$.

Example 1 (Cont) Consider the graph $\mathbf{G}_1$ and let $\epsilon = 0.75$. Hence, $v(c) = 1$ and $v(a) = v(b) = 0.75$. It is easy to check that $P_0(c) = P_1(c) = 1$, $P_0(a) = 0.75$, $P_1(a) = 0$ and $P_0(b) = 0.75$, $P_1(b) = -0.25$. Thus, $c \succeq_{\mathbf{G}_1}^P a \succeq_{\mathbf{G}_1}^P b$.

Ranking semantics provide several levels of acceptability. While this allows fine-grained comparisons of arguments, it may lead to move away from the essence of acceptability, which is predicting whether an argument can be accepted or not. Consider the graphs $\mathbf{G}_9$, $\mathbf{G}_{10}$, $\mathbf{G}_{11}$ from Table 1. The existing ranking semantics return only one comparison, namely $a \succeq_{\mathbf{G}_9} a$ and $a \succeq_{\mathbf{G}_{10}} a$. This means that they do not make any distinction between the two graphs $\mathbf{G}_9$, $\mathbf{G}_{10}$.

They do not state that a should be accepted in $\mathbf{G}_9$ and rejected in $\mathbf{G}_{10}$. They also declare a and b in $\mathbf{G}_{11}$ as equally acceptable, which is correct, but they do not declare both arguments as rejected. Burden-based semantics declares a as more acceptable than b in $\mathbf{G}_{12}$ but does not state that both are unacceptable. We will argue later that ranking semantics do not compare arguments with respect to their acceptability but rather regarding their *strengths*.

3 Strength, Acceptability, Coalitions

Argumentation is an approach for solving various AI problems including reasoning with conflicting information, classifying objects, making decisions, etc. It follows a four-step process: it 1) justifies claims by arguments, 2) identifies relations between pairs of arguments, 3) analyzes the arguments, and 4) identifies the output (eg. the conclusions to be drawn from a knowledge base, the class of an object, the prevailing opinion or the winner in a dialog, the option(s) to be chosen).

3.1 Three Goals = Three Concepts

The last two steps of an argumentation process are closely related. On the one hand, the output depends broadly on the results obtained in the analysis step. For instance, argument-based paraconsistent logics draw, from an inconsistent knowledge base, formulas that are supported by strong arguments, where the latter are identified in step 3. On the other hand, the nature of the output constraints the kind of analysis to be done in step 3. Below, we discuss three kinds analysis, but the list can largely be extended according to applications' needs.

Argument Strength

Recall that an argument is a set of premises that serve as reasons for accepting a claim. Its strength depends on the plausibility of the premises, the strength of the link between the premises and claim, and the prior acceptability of the claim. Attacks aim to highlight weaknesses in these three components of an argument. Hence, the less an argument is attacked, the stronger it is. Argument strength describes thus to what extent an argument is tarnished by attacks. It is a *many-valued* notion as the strength of an argument may range between completely strong and completely weak. Furthermore, an argument has a *unique* strength in a fixed graph. For instance, it is either strong or weak, it cannot be both. However, the value of an argument may be imprecise. Hence, an argument strength may be either *precise* (i.e., a single value) or *vague* (i.e., it belongs to an interval).

Characteristics: Uniqueness, Precise vs Vague.

Argument Acceptability

The idea is to predict whether an argument can be accepted so that its claim can safely be used for drawing other conclusions, making decisions, etc. Acceptability here is related to argument strength in the sense that only strong arguments can be accepted. It can however be defined in two ways: i) directly without computing strengths of arguments as done in (1), ii) it can be derived from arguments' strengths. For instance, one may accept any argument whose strength is greater than the strength of any of its attackers.

Formally, given a weighting semantics π , an argumentation graph $G = \langle \mathcal{A}, \mathcal{R} \rangle$, an argument $a \in \mathcal{A}$,

$$a \text{ is acceptable iff } \forall b \in \mathcal{R} \text{ such that } (b, a) \in \mathcal{R}, \text{Deg}_G^\pi(a) > \text{Deg}_G^\pi(b). \quad (2)$$

Another possibility consists of accepting all arguments having a strength above a certain threshold δ . Formally,

$$a \text{ is acceptable iff } \text{Deg}_G^\pi(a) > \delta. \quad (3)$$

Acceptability is a *binary* notion, a human agent either accepts or rejects an argument, she either relies or discards an argument's claim. However, one may also argue that among the acceptable arguments, some of them may be more acceptable than others. We have seen previously that skeptically accepted arguments are more acceptable than cleanly accepted ones. Finally, the status of an argument with respect to acceptability is *unique*. An argument cannot be both acceptable and unacceptable in a fixed argumentation graph.

Characteristics: Uniqueness, Binary vs Gradual.

Coalitions of Arguments

Unlike the two previous concepts, where the focus is on individual arguments, the concern here is the set of *viewpoints* expressed in an argumentation graph. Such a graph is seen as a power game between different competing *viewpoints*, and coalitions of arguments supporting prevailing viewpoints are looked for. Several coalitions may exist as several viewpoints can survive at the same time in a game. Furthermore, arguments may belong to zero, one, or several winning coalitions.

Characteristics: Multiple coalitions.

3.2 Which Concept to Choose?

The choice of the kind of analysis to perform depends on the problem to solve. For instance, in case of *multiple criteria decision making* (MCD), the main objective is to define mathematical models that are able to compare different alternatives on the basis of how they perform regarding a set of criteria. The more discriminating a model between alternatives, the more efficient it is. In argument-based MCD models, an argument in favor of an alternative expresses how the latter satisfies a criterion. Thus, it is not sufficient to identify which arguments are acceptable (as alternatives may all be supported by only acceptable arguments), their strengths are crucial for fine-grained comparisons of alternatives.

Consider the case of a judge who should make a decision knowing some conflicting arguments given by witnesses. The main goal of the judge is to know which arguments are acceptable, and can thus be wisely taken into account in his final judgment, and which ones should be ignored.

Consider now the case of online debate platforms (eg. Argüman), where people discuss about societal issues. One of the goals may be to identify prevailing viewpoints expressed by people on a given issue (eg. increasing taxes). Thus, one should consider the whole argumentation graph and analyze how arguments (opinions) are related to each other.

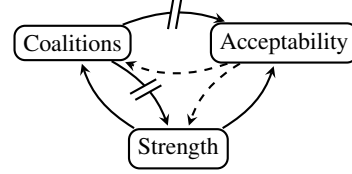


Figure 1: Links

Semantics	Goals
Weighting Semantics	Strength
Ranking Semantics	Strength
Extension Semantics	Acceptability

Table 3: Classification of existing semantics wrt goals

3.3 Links Between the Concepts

Acceptability can be defined from argument strength as shown previously. The probabilistic model defined in [Hunter, 2013] computes one set of acceptable arguments; it contains arguments having probability greater than 0.5. It is worth mentioning that an argumentation graph may not contain any acceptable argument if all its arguments are weak (ie., they have a low strength). Consider for instance the argumentation graph G_3 depicted in Figure 1. Assume that the weight of each argument belongs to the unit interval $[0, 1]$ and represents the certainty degree of the argument's premises. Those weights are too weak that even if c is not attacked, c may be considered as unacceptable. One may also reject the two arguments of the graph G_{11} since they are all self-attacked. From argument strength, one may also define *coalitions of strong arguments*, eg. conflict-free sets that contain arguments having a strength greater than a fixed threshold.

From argument acceptability, one may define a *binary* notion of strength: an argument is strong if it is acceptable and weak otherwise. A restricted notion of coalition can also be defined: a coalition contains non-conflicting acceptable arguments. Generally, one expects a single coalition since a rational agent may not accept an argument and its attackers.

Coalitions do not necessarily define acceptable arguments, i.e., elements of a coalition may not be acceptable. They do not either necessarily inform about strengths of arguments. For instance, one may consider that the argumentation graph G_3 has at least one coalition (defending a coherent viewpoint), the set $\{a, c\}$. However, both arguments may be seen as very weak, and cannot be accepted.

Figure 1 summarizes the links between the three concepts. Plain, Dashed, and Negated arrows stand respectively for existence, possible existence, non-existence of link.

4 Classification wrt Goals

The existing semantics can be classified as shown in Table 3.

4.1 Weighting Semantics

They evaluate the strength of every argument in a graph. This explains why they return a single vector of values (one value per argument) compared to multiple extensions (multiple labellings). Existing works on these semantics did not

tackle the question of acceptability (i.e., which arguments can be selected as acceptable). A notable exception is [Cayrol and Lagasque-Schiex, 2005], where the authors proposed the way (2) of selecting acceptable arguments (called well-defended arguments in that paper) using the values returned by a weighting semantics.

4.2 Ranking Semantics

These semantics are concerned with *argument strength*. They do not provide the list of acceptable arguments, but rather rank any pair of arguments even the weakest ones. For instance, existing ranking semantics declare the two arguments a and b of the graph G_{11} as equally strong, however they do not recommend them as rejected. In G_9 , none of existing ranking semantics declares a as acceptable.

An important question is then: *how do ranking semantics compare with weighting semantics?* First, the former rank arguments from the strongest to the weakest ones without necessarily computing their strengths. We call such semantics *pure* ranking semantics for distinguishing them from those that can be defined from weighting semantics. Note that every weighting semantics can give birth to a ranking one, it is sufficient to compare the values (strengths) of arguments. Second, unlike weighting semantics, they are unable to compare arguments of distinct graphs as shown in the two graphs G_9 and G_{10} . Existing ranking semantics do not distinguish between the two cases of a while h -Categorizer, for instance, assigns to a the value 1 in G_9 and 0.61 in G_{10} . Third, both classes of semantics obey to a similar list of principles. Indeed, the principles proposed in [Amgoud and Ben-Naim, 2013] for ranking semantics were later decomposed into more elementary ones in [Amgoud and Ben-Naim, 2016a] for weighting semantics. Fourth, unlike weighting semantics, it is hard to identify acceptable arguments from a ranking. Consider the graph G_{10} , any ranking semantics π would return $a \succeq_{G_{10}}^\pi a$. It is thus impossible to conclude that a should be rejected. Finally, while a ranking is meaningful, the values assigned by weighting semantics can hardly be interpreted.

Ranking/Weighting semantics are efficient in paraconsistent reasoning. [Amgoud and Ben-Naim, 2015] defined ranking logics using ranking semantics for the evaluation of arguments. Those logics are more productive (i.e., infer more conclusions) than those that use extension semantics.

4.3 Extension Semantics

They focus on acceptability of arguments, which is based on a particular principle (see (1)) and do not compute arguments' strengths as in (2) and (3). This explains why they violate most of the principles defined in [Amgoud and Ben-Naim, 2016a]. The values (skeptically accepted, credulously accepted, cleanly accepted, rejected) do not express strengths of arguments (i.e. $\text{Deg}^\pi(\cdot)$) but rather acceptability levels.

In some applications, one may need more information on arguments rather than just the fact of being acceptable/rejected. For instance, the graph G_2 has no stable extension, consequently all its arguments are rejected. However, one might argue that a is stronger than b , which is in turn stronger than c . Consider also the two graphs G_5 and G_6 . Under stable semantics, $\text{Ext}_{G_5}^s = \{\{a, c, d\}\}$ and

$\text{Ext}_{G_6}^s = \{\{a, c\}, \{a, d\}\}$. Note that a is acceptable in both graphs. However, one would argue that a should be stronger (then more acceptable) in G_5 since its two defenders in that graph are not attacked. Such information can be provided by an approach for acceptability that would be based on argument strength, like in (2) and (3). For instance, h -Categorizer assigns the value 0.75 to a in G_5 and 0.69 to a in G_6 .

The notion of acceptability as defined by extension semantics is based on some kind of coalitions of arguments, the one grounded on defence (see (1)). This explains partly the possibility of existence of several extensions. This *coalition-based acceptability* presents some weaknesses:

1) The acceptability status (eg., skeptically accepted) does not reflect argument strength, but rather the status of an argument in the coalition game. In the graph G_1 , the argument a is accepted while it was shown in different experimental studies [Rahwan *et al.*, 2010; Rosenfeld and Kraus, 2014; 2015] that humans do not follow the main principles behind extension semantics. In particular, those studies argued against the cognitive plausibility of reinstatement as the latter does not yield the full expected recovery of an attacked argument. The studies revealed that from the point of view of humans, a will lose adherence and may even be rejected.

2) Preferences between arguments or external weights of arguments (like certainty degrees of arguments' premises) are not properly considered by extension semantics. In the graph G_3 , all existing preference-based approaches for extension semantics ignore the weights of arguments (since attackers have greater weights than attackees) and return one stable extension $\{a, c\}$. Note that both arguments a and c are too weak that it is not wise to rely on their conclusions. Thus, the fact of belonging to an extension does not mean being acceptable.

3) The fact of being outside extensions does not mean either being weak. For instance, the graph G_2 has no stable extension while a can be considered as acceptable. Similarly, b may be considered as acceptable since its only attacker is the self-attacking argument c .

Acknowledgments

Support from the ANR-3IA Artificial and Natural Intelligence Toulouse Institute is gratefully acknowledged.

References

- [Amgoud and Ben-Naim, 2013] Leila Amgoud and Jonathan Ben-Naim. Ranking-based semantics for argumentation frameworks. In *Proc. of SUM*, pages 134–147, 2013.
- [Amgoud and Ben-Naim, 2015] Leila Amgoud and Jonathan Ben-Naim. Argumentation-based ranking logics. In *Proc. of AAMAS*, pages 1511–1519, 2015.
- [Amgoud and Ben-Naim, 2016a] Leila Amgoud and Jonathan Ben-Naim. Axiomatic foundations of acceptability semantics. In *KR*, pages 2–11, 2016.
- [Amgoud and Ben-Naim, 2016b] Leila Amgoud and Jonathan Ben-Naim. Evaluation of arguments from support relations: Axioms and semantics. In *Proc. of IJCAI*, pages 900–906, 2016.

- [Amgoud and Doder, 2019] Leila Amgoud and Dragan Doder. Acceptability semantics for weighted argumentation frameworks. In *Proc. of AAMAS*, pages 1270–1278, 2019.
- [Amgoud and Serrurier, 2008] Leila Amgoud and Mathieu Serrurier. Agents that argue and explain classifications. *JAAMAS*, 16(2):187–209, 2008.
- [Amgoud *et al.*, 2016] Leila Amgoud, Jonathan Ben-Naim, Dragan Doder, and Srdjan Vesic. Ranking arguments with compensation-based semantics. In *Proc. of KR*, pages 12–21, 2016.
- [Amgoud *et al.*, 2017] Leila Amgoud, Jonathan Ben-Naim, Dragan Doder, and Srdjan Vesic. Acceptability semantics for weighted argumentation frameworks. In *Proc. of IJCAI*, pages 56–62, 2017.
- [Atkinson *et al.*, 2017] Katie Atkinson, Pietro Baroni, Massimiliano Giacomin, Anthony Hunter, Henry Prakken, Chris Reed, Guillermo Ricardo Simari, Matthias Thimm, and Serena Villata. Towards artificial argumentation. *AI Magazine*, 38(3):25–36, 2017.
- [Baroni *et al.*, 2005] Pietro Baroni, Massimiliano Giacomin, and Giovanni Guida. SCC-recursiveness: a general schema for argumentation semantics. *Artificial Intelligence*, 168:162–210, 2005.
- [Baroni *et al.*, 2015] Pietro Baroni, Marco Romano, Francesca Toni, Marco Aurisicchio, and Giorgio Bertanza. Automatic evaluation of design alternatives with quantitative argumentation. *Argument & Computation*, 6(1):24–49, 2015.
- [Besnard and Hunter, 2001] Philippe Besnard and Anthony Hunter. A logic-based theory of deductive arguments. *Artificial Intelligence*, 128(1-2):203–235, 2001.
- [Bonzon *et al.*, 2016] Elise Bonzon, Jérôme Delobelle, Sébastien Konieczny, and Nicolas Maudet. Argumentation ranking semantics based on propagation. In *Proc. of COMMA*, pages 139–150, 2016.
- [Bonzon *et al.*, 2017] Elise Bonzon, Jérôme Delobelle, Sébastien Konieczny, and Nicolas Maudet. A parametrized ranking-based semantics for persuasion. In *Proc. of SUM*, pages 237–251, 2017.
- [Bonzon *et al.*, 2018] Elise Bonzon, Jérôme Delobelle, Sébastien Konieczny, and Nicolas Maudet. Combining extension-based semantics and ranking-based semantics for abstract argumentation. In *Proc. of KR*, pages 118–127, 2018.
- [Caminada, 2006] Martin Caminada. Semi-stable semantics. In *Proc. of COMMA*, pages 121–130, 2006.
- [Cayrol and Lagasquie-Schiex, 2005] Claudette Cayrol and Marie-Christine Lagasquie-Schiex. Graduality in argumentation. *JAIR*, 23:245–297, 2005.
- [da Costa Pereira *et al.*, 2011] Célia da Costa Pereira, Andrea Tettamanzi, and Serena Villata. Changing one’s mind: Erase or rewind? In *Proc. of IJCAI*, pages 164–171, 2011.
- [Dondio, 2018] Pierpaolo Dondio. Ranking semantics based on subgraphs analysis. In *Proc. of AAMAS*, pages 1132–1140, 2018.
- [Dung *et al.*, 2007] Phan Dung, Paolo Mancarella, and Francesca Toni. Computing ideal skeptical argumentation. *Artificial Intelligence*, 171:642–674, 2007.
- [Dung, 1995] Phan Minh Dung. On the Acceptability of Arguments and its Fundamental Role in Non-Monotonic Reasoning, Logic Programming and n-Person Games. *Artificial Intelligence*, 77:321–357, 1995.
- [Gabbay and Rodrigues, 2015] Dov Gabbay and Odinaldo Rodrigues. Equilibrium states in numerical argumentation networks. *Logica Universalis*, 9(4):411–473, 2015.
- [Hunter, 2013] Anthony Hunter. A probabilistic approach to modelling uncertain logical arguments. *I. Journal of Approximate Reasoning*, 54(1):47–81, 2013.
- [Leite and Martins, 2011] João Leite and João Martins. Social abstract argumentation. In *Proc. of IJCAI*, pages 2287–2292, 2011.
- [Lin and Shoham, 1989] Fangzhen Lin and Yoav Shoham. Argument systems - an uniform basis for non-monotonic reasoning. In *Proc. of KR*, pages 245 – 255, 1989.
- [Pu *et al.*, 2014] Fuan Pu, Jian Luo, Yulai Zhang, and Guiming Luo. Argument ranking with categoriser function. In *In Proc. of KSEM*, pages 290–301, 2014.
- [Rago *et al.*, 2016] Antonio Rago, Francesca Toni, Marco Aurisicchio, and Pietro Baroni. Discontinuity-free decision support with quantitative argumentation debates. In *Proc. of KR*, pages 63–73, 2016.
- [Rahwan *et al.*, 2010] Iyad Rahwan, Mohammed Iqbal Madakkattel, Jean Bonnefon, Ruqiyabi Naz Awan, and Sherief Abdallah. Behavioral experiments for assessing the abstract argumentation semantics of reinstatement. *Cognitive Science*, 34(8):1483–1502, 2010.
- [Rosenfeld and Kraus, 2014] Ariel Rosenfeld and Sarit Kraus. Argumentation theory in the field: An empirical study of fundamental notions. In *Proc. of the Workshop on Frontiers and Connections between Argumentation Theory and Natural Language Processing*, 2014.
- [Rosenfeld and Kraus, 2015] Ariel Rosenfeld and Sarit Kraus. Providing arguments in discussions based on the prediction of human argumentative behavior. In *Proc. of AAI*, pages 1320–1327, 2015.
- [Rosenfeld and Kraus, 2016] Ariel Rosenfeld and Sarit Kraus. Strategical argumentative agent for human persuasion. In *Proc. of ECAI*, pages 320–328, 2016.
- [Simari and Loui, 1992] Guillermo Ricardo Simari and Ronald Prescott Loui. A mathematical treatment of defeasible reasoning and its implementation. *Artificial Intelligence*, 53(2-3):125–157, 1992.
- [van der Torre and Vesic, 2017] Leon van der Torre and Srdjan Vesic. The principle-based approach to abstract argumentation semantics. *FLAP*, 4(8), 2017.
- [Zhong *et al.*, 2019] Qiaoting Zhong, Xiuyi Fan, Xudong Luo, and Francesca Toni. An explainable multi-attribute decision model based on argumentation. *Expert Systems with Applications*, 117:42–61, 2019.